

Análisis de la Investigación Cuantitativa

Métodos clásicos



Eduardo Jiménez Marqués



ÍNDICE

1.	CAPÍTULO 1 CONCEPTOS BÁSICOS.....	14
1.1	¿QUÉ ES LA INVESTIGACIÓN DE COMERCIAL?	14
1.2	¿QUÉ ES LA INVESTIGACIÓN CUANTITATIVA?	14
1.3	TÉCNICAS UTILIZADAS EN LA INVESTIGACIÓN CUANTITATIVA	15
1.3.1	LA OBSERVACIÓN	16
1.3.2	LA EXPERIMENTACIÓN	16
1.3.3	LA ENCUESTA ESTRUCTURADA	16
1.4	DETERMINACIÓN DE LA MUESTRA.....	16
1.4.1	DETERMINACIÓN DEL TAMAÑO DE LA MUESTRA	17
1.4.2	CASOS:.....	19
2.	CAPÍTULO 2 OBSERVACIÓN	21
2.1	DEFINICIÓN Y CONCEPTOS GENERALES	21
2.2	FASES DE LA INVESTIGACIÓN POR OBSERVACIÓN.....	22
2.3	TIPOS DE OBSERVACIÓN	24
2.3.1	EN FUNCIÓN DE LA INTERVENCIÓN DEL INVESTIGADOR	24
2.3.2	SEGÚN SE REALICE EN CIRCUNSTANCIAS NATURALES O ARTIFICIALES ...	25
2.3.3	SEGÚN LA PARTICIPACIÓN DE LA MUESTRA	25
2.3.4	DE ACUERDO CON SU ESTRUCTURACIÓN	25
2.3.5	EN FUNCIÓN DE LA FORMA	26

Análisis de la Investigación Cuantitativa

2.3.6	DE CONFORMIDAD CON EL PROCEDIMIENTO	26
2.4	ELABORACIÓN DE UN CÓDIGO DE OBSERVACIÓN	26
2.5	MEDIDAS EN LA INVESTIGACIÓN POR OBSERVACIÓN	27
2.6	TÉCNICAS DE MUESTREO Y OBSERVACIÓN	28
2.7	EVALUACIÓN DE LA OBSERVACIÓN.....	29
2.7.1	CASO PRÁCTICO	31
2.8	VALIDEZ	32
2.8.1	VALIDEZ DE CONTENIDO	33
2.8.2	VALIDEZ DE CONSTRUCTO.....	33
2.8.3	VALIDEZ ORIENTADA AL CRITERIO	33
2.9	ANÁLISIS DE DATOS	33
2.10	FUENTES DE ERROR EN LA OBSERVACIÓN.....	33
2.11	VENTAJAS E INCONVENIENTES DE LA INVESTIGACIÓN POR OBSERVACIÓN.....	34
3.	CAPÍTULO 3 EXPERIMENTACIÓN	35
3.1	INTRODUCCIÓN.....	35
3.2	CONCEPTOS GENERALES.....	35
3.3	PRINCIPALES APLICACIONES	37

Análisis de la Investigación Cuantitativa

3.4	METODOLOGÍA	37
3.5	INDICADORES DE VALIDEZ DE UN EXPERIMENTO	38
3.6	ANÁLISIS ESTADÍSTICO DE LOS DISEÑOS EXPERIMENTALES.....	38
3.6.1	MÉTODO ANOVA TRADICIONAL.....	39
3.7	TIPOS DE EXPERIMENTOS HABITUALMENTE UTILIZADOS EN INVESTIGACIÓN COMERCIAL.....	40
3.7.1	EXPERIMENTACIÓN AL AZAR.....	41
3.7.2	EXPERIMENTACIÓN EN BLOQUES ALEATORIOS	41
3.7.3	EXPERIMENTO DE CUADRADO LATINO	41
3.7.4	EXPERIMENTO CON INTERCAMBIO	41
3.7.5	EXPERIMENTO FACTORIAL.....	41
3.8	LIMITACIONES DE LA EXPERIMENTACIÓN COMERCIAL.....	42
3.9	CASO PRÁCTICO: EXPERIMENTO DE BLOQUE ALEATORIO.....	42
3.9.1	EXPERIMENTO DE BLOQUE ALEATORIO (SPSS)	45
3.10	TABLA ESTADÍSTICA: DISTRIBUCIÓN DE LA F.....	46
3.11	MÉTODOS DE SIMULACIÓN.....	46
3.12	CONCEPTO.....	46
3.13	CLASIFICACIÓN.....	48
3.14	VENTAJAS Y LIMITACIONES.....	48
4.	CAPÍTULO 4 LA ENCUESTA ESTRUCTURADA	49

4.1	LA ENCUESTA. TIPOS. DATOS QUE APORTA. VENTAJAS E INCONVENIENTES.....	49
4.1.1	TIPOS DE ENCUESTA ESTRUCTURADA.....	51
4.1.2	INFORMACIÓN QUE APORTAN LAS ENCUESTAS.....	52
4.1.3	VENTAJAS E INCONVENIENTES	52
4.2	CARACTERÍSTICAS DE LAS ENCUESTAS ESTRUCTURADAS.....	53
4.3	TIPOS DE ENCUESTA ESTRUCTURADA SEGÚN EL PROCEDIMIENTO EMPLEADO PARA OBTENER LA INFORMACIÓN	53
4.4	ENCUESTA PERSONAL VIS A VIS.....	53
4.4.1	ENTREVISTAS POR DETENCIÓN	55
4.4.2	ESTUDIOS CON INVITADOS A UNA LOCALIZACIÓN CENTRAL	56
4.5	ENCUESTA TELEFÓNICA.....	56
4.6	ENCUESTA POR CORREO.....	58
4.6.1	MIXTAS COMBINACIONES TELEFÓNICA, CORREO Y PERSONAL	61
4.7	ENCUESTAS ESTRUCTURADAS POR SUSCRIPCIÓN.....	62
4.8	ENCUESTA SECTORIAL DE BIENES DE CONSUMO DURADERO...	62
4.9	ENCUESTA ÓMNIBUS.....	63
4.10	EL PANEL.....	64
4.10.1	PANEL DE CONSUMIDORES.....	65
4.10.2	PANEL DE HOGARES	66
4.10.3	DUSTBIN - CHEK.....	68

Análisis de la Investigación Cuantitativa

4.10.4	PANEL DE AUDIÓMETROS (T.V.)	68
4.10.5	PANEL DE DETALLISTAS.....	69
4.11	DIFERENCIAS ENTRE LAS DISTINTAS MODALIDADES DE ENCUESTA ESTRUCTURADA.....	71
5.	CAPÍTULO 5 INTRODUCCIÓN AL ANÁLISIS.....	72
5.1	INTRODUCCIÓN.....	72
5.2	FASES DEL PROCESO DE ANÁLISIS DE LOS DATOS	72
5.2.1	REVISIÓN DEL TRABAJO DE CAMPO Y DE LOS CUESTIONARIOS.....	72
5.2.2	CODIFICACIÓN Y TABULACIÓN	73
5.2.3	ANÁLISIS DE CADA CUESTIÓN O ÍTEM	75
5.2.4	ANÁLISIS DE LOS ÍTEMS POR SUBGRUPOS.	75
5.2.5	ESTUDIO DE LAS RELACIONES ENTRE PARES DE PREGUNTAS	76
5.2.6	ESTUDIO DE LAS RELACIONES ENTRE TODAS LAS PREGUNTAS:	76
5.2.7	RESULTADOS. CONCLUSIONES. INFORME	76
5.3	TABLAS DE DATOS	77
5.4	TIPOS DE TABLAS.....	78
5.4.1	TABLAS CUANTITATIVAS	78
5.4.2	TABLAS DE DATOS ORDINALES Y PREFERENCIAS	78
5.4.3	TABLAS BINARIAS	79
5.4.4	TABLAS DE MODALIDADES	79
5.4.5	TABLAS DISYUNTIVAS COMPLETAS	80
5.4.6	TABLAS DE PROXIMIDADES Y DISTANCIAS.....	80
5.4.7	TABLAS DE SERIES TEMPORALES	80

5.4.8	TABLAS MIXTAS O HETEROGÉNEAS	81
6.	CAPÍTULO 6 ANÁLISIS DE LOS DATOS	82
6.1	MÉTODOS ESTADÍSTICOS CLÁSICOS	83
6.1.1	CONCEPTOS BÁSICOS	83
6.1.1.1	TIPOS DE ESTADÍSTICA.....	83
6.1.1.2	MEDIR	83
6.1.1.3	ATRIBUTO.....	83
6.1.1.4	VARIABLE	84
6.1.2	ESCALAS DE MEDIDA	84
6.1.2.1	ESCALA NOMINAL.....	84
6.1.2.2	ESCALA ORDINAL	85
6.1.2.3	ESCALA DE INTERVALO	86
6.1.2.4	ESCALA DE RAZÓN O PROPORCIÓN	86
6.1.3	ELEMENTOS DE LA ESTADÍSTICA INFERENCIAL	87
6.2	ANÁLISIS UNIVARIABLE.....	88
6.2.1	DESCRIPCIÓN DE LOS DATOS	88
6.2.2	FRECUENCIA	88
6.2.3	MEDIDAS DE TENDENCIA CENTRAL.....	89
6.2.4	MEDIDAS DE DISPERSIÓN	90
6.2.4.1	AMPLITUD O RANGO	91
6.2.4.2	RECORRIDO INTERCUARTÍLICO	92
6.2.4.3	LA DESVIACIÓN INTERCUARTIL	92
6.2.4.4	DESVIACIÓN MEDIA	93
6.2.4.5	VARIANZA Y DESVIACIÓN TÍPICA O ESTÁNDAR	93
6.2.4.6	COEFICIENTE DE VARIACIÓN	94
6.2.5	MEDIDAS RELATIVAS A LA FORMA DE LA DISTRIBUCIÓN.....	95

6.2.5.1	DISTRIBUCIÓN	95
6.2.5.2	COEFICIENTE DE SESGO O ASIMETRÍA (SKEWNESS)	95
6.2.5.3	COEFICIENTE DE CURTOSIS O APUNTAMIENTO	96
6.2.6	¿CÓMO REALIZAR INFERENCIAS?	98
7.	CAPÍTULO 7 CONTRASTE DE HIPÓTESIS.....	99
7.1	CONCEPTOS GENERALES.....	99
7.1.1	HIPÓTESIS	99
7.1.2	PRUEBAS DE HIPÓTESIS	99
7.1.2.1	HIPÓTESIS NULA	100
7.1.2.2	HIPÓTESIS ALTERNATIVA.....	100
7.1.3	CARACTERÍSTICAS Y METODOLOGÍA	100
7.2	TEST DE HIPÓTESIS.....	104
7.2.1	OBJETIVO	104
7.3	METODOLOGÍA DEL TEST DE HIPÓTESIS	104
7.3.1	FORMULACIÓN DE LAS HIPÓTESIS.....	105
7.3.2	ELECCIÓN DEL NIVEL DE SIGNIFICACIÓN	106
7.3.3	ELECCIÓN DEL TEST	108
7.3.4	INTERPRETACIÓN DE LA PRUEBA	111
7.4	TIPOS DE TEST DE HIPÓTESIS	112
7.5	BREVE DESCRIPCIÓN DE LOS TEST.....	114
7.5.1	CONTRASTES PARA UNA MUESTRA	114
7.5.2	CONTRASTES PARA DOS MUESTRAS INDEPENDIENTES	115
7.5.3	CONTRASTES PARA DOS MUESTRAS RELACIONADAS	116

8.	CAPITULO 8: ANÁLISIS BIVARIABLE.....	117
8.1	RELACIÓN ENTRE DOS VARIABLES CUALITATIVAS	117
8.1.1	TABLAS DE CONTINGENCIA.....	117
8.1.1.1	TABLA DE CONTINGENCIA : SEXO Y CONOCIMIENTO DE INFORMÁTICA.....	118
8.1.1.2	SIGNIFICADO DE LOS ELEMENTOS QUE COMPONEN LA TABLA	119
8.1.1.3	ANÁLISIS DE INDEPENDENCIA ENTRE LAS DOS VARIABLES. ESTADÍSTICO χ^2 (CHI CUADRADO).....	120
8.1.2	COEFICIENTE V DE CRAMER	123
8.1.3	CORRELACIÓN DE RANGOS DE SPEARMAN	123
8.2	MÉTODOS DE MEDICIÓN EN EL ANÁLISIS ENTRE DOS VARIABLES CUANTITATIVAS	124
8.2.1	CORRELACIÓN	124
8.2.1.1	EJEMPLO	126
8.2.1.2	TABLA COEFICIENTE DE CORRELACIÓN R DE PEARSON.....	127
8.2.2	REGRESIÓN SIMPLE	128
8.2.2.1	OBJETIVOS	129
8.2.2.2	EJEMPLO REGRESIÓN LINEAL	130
8.2.3	COEFICIENTE ALFA DE CRONBACH	131
8.3	RELACIÓN ENTRE UNA VARIABLE CUANTITATIVA Y OTRA CUALITATIVA.....	132
8.3.1	ANÁLISIS DE LA VARIANZA	132
8.3.1.1	MÉTODO ANOVA TRADICIONAL UNIDIRECCIONAL.....	133
8.3.1.2	TABLA ESTADÍSTICA: DISTRIBUCIÓN DE LA F (NIVEL DE CONFIANZA 95%)	135

8.3.1.3	EJEMPLO DE ANOVA UNIDIRECCIONAL	137
8.3.2	TEST T DE MEDIAS.....	140
9.	CAPITULO 9 ANÁLISIS DE MUESTRAS PEQUEÑAS.....	141
9.1	INTRODUCCIÓN.....	141
9.2	DISTRIBUCIÓN “T” DE STUDENT.....	141
9.2.1	FIABILIDAD DE UN ESTADÍSTICO	143
9.2.2	SIGNIFICACIÓN DE LA MEDIA DE MUESTRAS PEQUEÑAS Y SU FIABILIDAD	144
9.2.2.1	CASO PRÁCTICO	144
9.2.3	SIGNIFICACIÓN DE LA DIFERENCIA DE MEDIAS EN MUESTRAS PEQUEÑAS INDEPENDIENTES	145
9.2.3.1	CASO PRÁCTICO	145
9.2.4	SIGNIFICACIÓN DE LA DIFERENCIA DE MEDIAS DE MUESTRAS PEQUEÑAS RELACIONADAS	146
9.2.4.1	CASO PRÁCTICO	146
9.2.5	SIGNIFICACIÓN DE LA DIFERENCIA DE DESVIACIONES TÍPICAS DE MUESTRAS RELACIONADAS	148
9.2.5.1	CASO PRÁCTICO	148
9.3	SIGNIFICACIÓN Y FIABILIDAD DEL COEFICIENTE DE CORRELACIÓN R DE PEARSON PARA MUESTRAS PEQUEÑAS	149
9.3.1	CASO PRÁCTICO	149
9.3.2	SIGNIFICACIÓN DE LA DIFERENCIA ENTRE COEFICIENTES DE CORRELACIÓN OBTENIDOS EN MUESTRAS RELACIONADAS.....	150
9.3.2.1	CASO PRÁCTICO	150
9.4	DISTRIBUCIÓN CHI CUADRADO (χ^2)	152

Análisis de la Investigación Cuantitativa

9.4.1	CASO PRÁCTICO 1	152
9.4.2	CASO PRÁCTICO 2	153
9.5	DISTRIBUCIÓN F DE FISHER.....	155
9.5.1	CASO PRÁCTICO	156
10.	CAPÍTULO 10 TEST PARAMÉTRICOS.....	158
10.1	INTRODUCCIÓN	158
10.2	TIPOS DE TEST PARAMÉTRICOS.....	158
10.2.1	CONTRASTES PARA UNA MUESTRA	158
10.2.2	CONTRASTES PARA DOS MUESTRAS INDEPENDIENTES	159
10.2.3	CONTRASTES PARA DOS MUESTRAS RELACIONADAS	159
10.2.4	PRUEBAS MÁS UTILIZADAS	159
10.3	TEST DE LA MEDIA.....	159
10.3.1	PRUEBA Z	159
10.3.2	PRUEBA T	160
10.3.3	PRUEBA DE DIFERENCIA DE MEDIAS INDEPENDIENTES.....	165
10.3.3.1	CON VARIANZA CONOCIDA	165
10.3.3.2	CON VARIANZAS DESCONOCIDAS.....	166
10.4	TEST DE PROPORCIONES.....	168
10.5	TEST DE SIGNIFICACIÓN PARA POBLACIONES INFINITAS Y QUE SIGUEN LA DISTRIBUCIÓN NORMAL	168
10.5.1	TEST DE SIGNIFICACIÓN PARA DIFERENCIAS DE PROPORCIONES INDEPENDIENTES	169

Análisis de la Investigación Cuantitativa

10.5.2	TEST DE SIGNIFICACIÓN PARA DIFERENCIAS DE PROPORCIONES NO INDEPENDIENTES	170
10.5.2.1	PORCENTAJES EXCLUYENTES	170
10.5.2.2	PORCENTAJES SOLAPADOS.....	171
10.6	TEST PARA MUESTRAS RELACIONADAS.....	172
11.	CAPÍTULO 11 TEST NO PARAMÉTRICOS (I).....	173
11.1	INTRODUCCIÓN	173
11.2	CLASIFICACIÓN DE LOS TEST NO PARAMÉTRICOS.....	173
11.3	BREVE DESCRIPCIÓN DE DIFERENTES TEST NO PARAMÉTRICOS	175
11.3.1	INTRODUCCIÓN	175
11.3.2	UNA MUESTRA MEDIDA UNA SOLA VEZ	175
11.4	TEST DE LA CHI CUADRADO	176
11.4.1	CASO PRÁCTICO	177
11.5	PRUEBA DE LA BINOMIAL.....	180
11.5.1	EJEMPLO:.....	181
11.6	PRUEBA BINOMIAL PARA MUESTRAS PEQUEÑAS	182
11.7	PRUEBA BINOMIAL PARA MUESTRAS GRANDES.....	182
11.7.1	CASO PRÁCTICO	184
11.8	TEST DE KOLMOGOROV – SMIRNOV (KS).....	187
11.8.1	CASO PRÁCTICO	188

11.9	CASO DE UNA MUESTRA MEDIDA DOS VECES.....	191
11.9.1	TEST DE MCNEMAR	191
11.9.1.1	CASO PRÁCTICO	193
11.9.2	TEST DE LOS SIGNOS	196
11.9.2.1	APLICACIÓN EN MUESTRAS PEQUEÑAS.	197
11.9.2.2	APLICACIÓN EN MUESTRAS GRANDES.....	197
11.9.3	TEST DE RANGOS ASIGNADOS DE WILCOXON	200
11.9.3.1	APLICACIÓN EN MUESTRAS PEQUEÑAS	202
11.9.3.2	MUESTRAS GRANDES	202
11.10	CASO DE DOS MUESTRAS INDEPENDIENTES.	206
11.10.1	INTRODUCCIÓN	206
11.11	CASO DE K MUESTRAS RELACIONADAS.....	206
11.11.1	INTRODUCCIÓN	206
11.12	CASO DE “K” MUESTRAS INDEPENDIENTES.....	207
11.12.1	INTRODUCCIÓN	207
12.	BIBLIOGRAFÍA RECOMENDADA.....	208

1. CONCEPTOS BÁSICOS

1.1 ¿QUÉ ES LA INVESTIGACIÓN DE COMERCIAL?

Podemos afirmar que la Investigación Comercial o de Marketing consiste en sentir y escuchar al consumidor. Prácticamente todas las organizaciones buscan información que les permita saber qué es lo que la gente quiere y por qué lo quiere, siendo los análisis más profundos los que determinan que el consumidor sea el que mejor lo sabe.

La base de la Investigación de Marketing descansa en el sentido común. No hay nada de paranormal ni en sus objetivos ni en sus métodos.

Cualquier persona que va a tomar una decisión reflexiva, realiza una investigación con el fin de reducir al mínimo el riesgo de equivocarse y de no obtener la satisfacción que busca. Algunos ejemplos cotidianos son: la compra de un coche, un piso, un electrodoméstico, etc. En todos los casos, seguimos un proceso de investigación previo a la toma de decisión, con el fin de buscar la alternativa mas adecuada a nuestras necesidades.

Podemos definir la Investigación de Marketing como el conjunto de técnicas enfocadas a obtener información objetiva sobre el mercado, con el fin de facilitar la toma de decisiones con el mínimo de incertidumbre (que resulten lo más acertadas posibleS y con el mínimo riesgo).

1.2 ¿QUÉ ES LA INVESTIGACIÓN CUANTITATIVA?

La Investigación Cuantitativa de Marketing comprende un conjunto de técnicas destinadas a obtener información sobre los consumidores, centrada en su comportamiento externo, siendo su resultado siempre cuantificable.

Para obtener la información, se acostumbra a escoger elementos representativos de la población que se utiliza en la investigación. Al grupo de unidades representativas de la población que se quiere estudiar, y que se acostumbra a seleccionar según criterios estadísticos, se le denomina muestra, apareciendo siempre un margen de error que se asume en toda investigación muestral.

La investigación cuantitativa aporta información sobre aspectos cuantificables del mercado. Las técnicas de Investigación Cuantitativas responden a las preguntas ¿qué?, ¿cuánto?, ¿dónde?, ¿quién?, ¿cuándo?, ¿cómo?, pero no a la pregunta ¿por qué?

Los principales aspectos que se recogen son:

- Nivel de consumo de los diferentes productos o marcas
- Lugares de compra o puntos de venta donde los consumidores efectúan sus compras
- Número de clientes que compran en un determinado establecimiento
- Precios, niveles de precios en función de diversos parámetros
- Niveles de existencias en los canales de distribución
- Estudios de consumo por hábitat, comunidad autónoma (CCAA), región ...
- Estudios de resultados de la competencia
- Frecuencias de compra
- Volúmenes de ventas y compras por diversos conceptos

La Investigación Cuantitativa se caracteriza porque la información procede de los niveles superficiales de la personalidad de los consumidores, mientras que la Investigación Cualitativa obtiene su información investigando el interior de la personalidad de los individuos que integran la muestra.

En la praxis, se suelen emplear en muchos casos los dos tipos de investigación, por eso se suele hablar de técnicas mixtas.

Como resumen, podemos decir que la Investigación Cuantitativa trata del análisis de diferentes aspectos mensurables del mercado, obteniéndose normalmente la información a través de muestras representativas del colectivo objeto de estudio, extrapolándose los resultados obtenidos a toda la población con un determinado nivel de confianza, probabilidad y error estadístico predeterminado.

1.3 TÉCNICAS UTILIZADAS EN LA INVESTIGACIÓN CUANTITATIVA

Las técnicas utilizadas en la investigación cuantitativa son:

La observación.

La experimentación

La encuesta estructurada

Seguidamente, y de una forma sucinta, vamos a indicar en qué consiste cada una de estas técnicas.

1.3.1 LA OBSERVACIÓN

Se trata de una técnica de recogida de la información que, básicamente, consiste en observar y registrar las acciones, comportamientos, hechos y actuaciones de los individuos objeto de estudio, tal y como lo hacen habitualmente, sin que conozcan que están siendo estudiados.

1.3.2 LA EXPERIMENTACIÓN

Es un procedimiento de contrastar, así como de descubrir hipótesis. Mediante la experimentación se puede analizar el efecto que una variable independiente produce sobre otra variable dependiente. Para ello, se debe neutralizar y controlar la influencia que otros factores pudiesen ejercer sobre la variable objeto de estudio.

1.3.3 LA ENCUESTA ESTRUCTURADA

La encuesta estructurada es un método de obtención de datos mediante entrevista individual, en la que el entrevistado proporciona información de forma voluntaria y consciente, como respuesta a una serie de preguntas planteadas en el cuestionario.

1.4 DETERMINACIÓN DE LA MUESTRA

En la investigación comercial la población objeto de estudio suele comprender un numeroso grupo de elementos o individuos. Ante la imposibilidad de estudiar a todos los componentes de esa población o universo (N), lo que se hace es estudiar un número limitado de elementos, los cuales se deberán seleccionar adecuadamente.

Para dar respuesta a este problema existen técnicas derivadas de las matemáticas y especialmente en la estadística, las denominadas técnicas de muestreo, que nos permiten seleccionar adecuadamente a los elementos de la población para que conformen la muestra (n).

El número de individuos o elementos a seleccionar constituye el tamaño de la muestra.

La muestra debe cumplir las siguientes condiciones:

- Comprender parte del universo y no la totalidad del mismo
- Debe ser representativa del universo del que se ha extraído

No deben existir desviaciones en la elección de los elementos que formarán parte de la muestra.

Trabajar con muestras nos lleva a que el valor obtenido, también llamado valor estimado, pueda no coincidir exactamente con el valor real de la población. A la diferencia existente entre ambos valores se le denomina error muestral.

Para disminuir este error absoluto deberemos trabajar con muestras suficientemente grandes y elegir sus elementos con el método de muestreo adecuado.

1.4.1 DETERMINACIÓN DEL TAMAÑO DE LA MUESTRA

En la determinación del tamaño de una muestra intervienen los siguientes aspectos:

- Error de muestreo: Es decir, cual será la diferencia máxima que se admitirá entre los valores estimados y los reales o parámetros de la población. Este valor se debe fijar a priori.

Ejemplo si al determinar el consumo medio por habitante, admitimos un error de “e” unidades, si el valor medio del universo es μ y el obtenido mediante la muestra es **m**, el error absoluto **e** vendrá dado por la siguiente fórmula:

$$e = \mu - m$$

- El nivel o intervalo de confianza: Representa la probabilidad con que se desea que el estudio cumpla que la diferencia entre el valor estimado (m) y el valor real (μ) esté comprendida en los márgenes del error absoluto

$$Pr ob[|m - \mu| \leq e] = P_k$$

P_k dependerá de la distribución estadística del estimador.

El error muestral es la desviación típica del estimador.

En el caso de la media, viene dado por la siguiente fórmula o ecuación

$$s_m = \sqrt{\frac{N-n}{N-1} \frac{S^2}{n}}$$

Siendo N el universo o población, n el tamaño de la muestra y S^2 la cuasivarianza poblacional.

$$S^2 = \frac{\sum_{i=1}^N (x_i - m)^2}{N-1}$$

Para una mejor comprensión de la relación existente entre estos tres conceptos y el tamaño de la muestra, vamos a considerar que queremos estimar el valor medio y que la

distribución de valores que se obtendrá como resultado de la muestra, de la variable cuya media queremos estimar, sigue la distribución normal.

Una propiedad de la curva normal es aquella que liga la media de todos los valores de la distribución a la desviación típica de los valores y, en consecuencia, el valor de la razón crítica (K o Z) con la probabilidad (P).

De esta forma el intervalo de confianza queda definido por la siguiente fórmula:

$$\text{Pr ob.} [m - m] \leq K s_m = P_K$$

Si tenemos en cuenta que :

$$\text{Pr ob} [|m - m| \leq e] = P_K$$

se deduce fácilmente que el error absoluto, en el caso más desfavorable, será igual a K veces la desviación típica del estimador (normalmente la media o la proporción). Para el caso de la media será:

$$e = K \sigma_m$$

Y para la proporción será

$$e = K \sigma_p$$

Por consiguiente obtenemos que:

$$e = K \sqrt{\frac{N-n}{N-1} \cdot \frac{S^2}{n}}$$

Por tanto el tamaño de la muestra n, será

$$n = \frac{K^2 N S^2}{N e^2 + K^2 S^2}$$

En el caso de la proporción el tamaño de la muestra vendrá dado por

$$n = \frac{K^2 N p q}{e^2 (N-1) + K^2 p q}$$

En la Investigación comercial se suele trabajar en niveles de confianza del 95% en cuyo caso $K = 1.96$, y niveles de dos sigmas o 95.5% que equivale a $K = 2$.

También se trabaja con poblaciones o universos denominados infinitos, es decir de más de 100.000 elementos. En este caso, y a partir del valor 100.000, a igualdad de intervalo o nivel de confianza y de error absoluto, el tamaño del universo ya no influye en el de la muestra. Es decir, la relación

$$\frac{N-n}{N-1} \text{ tiende a } 1$$

Con lo que las formulas para el tamaño de la muestra en el caso de poblaciones infinitas vienen dadas por las siguientes expresiones, según nos refiramos a medias o proporciones

$$n = \frac{K^2 S^2}{e^2} \quad \text{y} \quad n = \frac{K^2 p q}{e^2}$$

La más utilizada habitualmente es la correspondiente a la proporción que ante la falta de información se toma la situación más desfavorable, es decir $p = q = 0.5$

1.4.2 CASOS:

Ejemplo 1. Una empresa de productos de gran consumo quiere realizar un estudio en la Unión Europea, para ello quiere tomar una muestra representativa de hogares de la UE. Las condiciones para la realización del trabajo son: el error máximo admitido será del 2%, para un nivel de confianza del 95% (Dos sigmas $K = 2$), y $p = q = 0.5$. Determina el tamaño de la muestra.

RESPUESTA

Como se trata de un universo infinito (N es mayor a 100.000 Hogares), la fórmula a utilizar será

$$n = \frac{k^2 pq}{e^2} \text{ sustituyendo obtenemos}$$
$$n = \frac{2^2 50 50}{2^2} = 2.500 \text{ hogares}$$

Ejemplo 2. Un fabricante de componentes industriales quiere realizar un estudio sobre empresas de un sector industrial susceptibles de comprar un producto novedoso. Para ello se plantea la realización de una encuesta estructurada a los jefes de compras de las diferentes empresas del sector. Teniendo en cuenta que el sector industrial objeto de estudio está conformado por 50.000 empresas y que la encuesta se va a realizar con un error máximo del 2%, y se conoce por otros estudios que $p = 80\%$. Determina la muestra necesaria para un nivel de confianza del 95% ($k=2$).

RESPUESTA

Se trata de un universo finito $N = 50.000$ empresas, como $p = 80\%$ y q será

$$q = 100 - 80 = 20\%$$

Investigación Comercial

Análisis de la Investigación Cuantitativa

El tamaño de la muestra se determina de acuerdo con la siguiente fórmula

$$n = \frac{k^2 pq N}{(N-1) + k^2 pq} \text{ sustituyendo obtenemos}$$
$$n = \frac{2^2 80 20 50.000}{(50.000-1)2^2 + 2^2 80 20} = 1550'4 = 1551 \text{ empresas}$$

Ejemplo 3. Se ha realizado un estudio de mercado en el sector comercio y queremos conocer si la muestra utilizada es significativa de esa población.

Se ha trabajado con una muestra de 280 tiendas, lo que presupone que el error del 5'78%.

Con el objetivo de extrapolar los resultados obtenidos a la población, se quiere demostrar que la muestra obtenida es representativa de la población objeto de estudio con el margen de error especificado (5'78%).

En la tabla siguiente figura la distribución de los comercios en función de su actividad principal en el universo o población y los resultados obtenidos en la muestra

	población	Muestra	
actividad	porcentaje	porcentaje	diferencia
Alimentación	38'7	37'9	- 0'8
Textil	19'5	19'6	0'1
Ferretería	6'3	7'1	0'8
Hogar	12'3	14'6	2'3
Vehículos	0'8	0'4	- 0'4
Otros	22'4	20'4	- 2
TOTAL	100'0	100'0	

CONCLUSIÓN

Comparando los estadísticos obtenidos en la muestra con los de la población, se aprecia que en ningún caso la diferencia entre ambos supera el máximo error especificado (5'78%), garantizando por consiguiente la representatividad de la muestra y la extrapolación de los resultados a la población.

2. OBSERVACIÓN

2.1 DEFINICIÓN Y CONCEPTOS GENERALES

En el área del marketing, la observación es una técnica de recogida de información que consiste básicamente en observar y recoger las actuaciones, comportamientos y hechos de las personas, tal y como los realizan habitualmente.

Ante un problema, las personas buscan asesoramiento en diferentes ámbitos, normalmente pidiendo consejo a otras personas que puedan saber más sobre aquello que les preocupa. A través de la tradición oral y escrita, la humanidad ha ido acumulando información acerca de sí mismos y de la naturaleza. Este cúmulo de experiencias acumuladas a lo largo de la historia de cada cultura es lo que se denomina conocimiento.

Para adquirir el conocimiento, el hombre utiliza diversos procedimientos, siendo el más antiguo y a la vez el más moderno el de la observación; podemos afirmar que la ciencia comienza con la observación. No se trata de una observación superficial, sino de una observación científica. Se trata de conseguir un procedimiento o método que permita la replicabilidad, es decir un método que seguido por un investigador haga posible que cualquier otro colega, cuando estudie el mismo fenómeno, obtenga idénticos resultados.

Los métodos utilizados son:

método inductivo método deductivo y método hipotético deductivo

Método inductivo. El método inductivo se ha desarrollado desde la postura que valora la experiencia como punto de partida para la generación del conocimiento, es decir, se parte de la observación de la realidad para, mediante la generalización de la observación, formular la ley.

Método deductivo. El método deductivo parte de la ley general, a la que se llega mediante la razón, y de ella se deducen consecuencias lógicas aplicables a la realidad.

Método hipotético deductivo. El método hipotético deductivo es una combinación de los dos anteriores. Trata de enfatizar el hecho de que el proceso de adquisición de nuevos conocimientos actúa de forma tal que el investigador necesita tanto ir de los datos a la teoría, como de ésta a los datos.

Desde una teoría se deducen consecuencias contrastables en la realidad; para ello se realizan una serie de observaciones que sirven para corroborar o modificar lo deducido desde la teoría.

En el caso de que no exista una teoría, se puede empezar realizando una observación a partir de la cual se hará la generalización.

La observación como método científico

Si observar es advertir los hechos como espontáneamente se presentan y consignarlos por escrito, en primer lugar se perciben tales hechos, los cuales, después, se expresan mediante palabras, signos u otras manifestaciones. Precisamente, el fundamento de la observación científica reside en la comprobación del fenómeno que se tiene frente a la vista, con la única preocupación de evitar y prever los errores de observación que podrían alterar la percepción de un fenómeno o la correcta expresión de éste.

El observador se diferencia del testigo ordinario de los hechos en que este último no intenta llegar a un diagnóstico de los mismos y en que, además, la mayor parte de los sucesos le pasan desapercibidos.

La observación se convierte en técnica científica en la medida en que:

- Sirve a un objetivo ya formulado de investigación
- Es planificada sistemáticamente
- Es controlada y relacionada con proposiciones más generales en vez de ser presentada como una serie de curiosidades interesantes
- Está sujeta a comprobaciones de validez y fiabilidad

Esto significa que deben formularse unas hipótesis a partir de una exploración empírica de las situaciones que se tratan de esclarecer. Seguidamente, se verifican las hipótesis, confrontándolas con el mayor número posible de hechos revelados por investigaciones, llegando de esta forma a un diagnóstico válido de la situación y a la elaboración de una teoría aplicable a la generalidad de los fenómenos del mismo tipo.

2.2 FASES DE LA INVESTIGACIÓN POR OBSERVACIÓN

La aplicación de la investigación por observación no se puede realizar de cualquier manera, sino que deberemos plantearnos y decidir qué, cómo y cuándo observar.

Análisis de la Investigación Cuantitativa

Para que la observación sea válida en marketing, se requiere que sea sistemática, es decir, deberemos realizarla de tal manera que dé lugar a datos susceptibles de ser reproducidos por cualquier otro investigador.

En primer lugar, hemos de partir del principio de que, para poder observar algo, primero habrá que tener en cuenta qué es lo que queremos conocer, es decir, los objetivos de la investigación.

El problema de qué observar viene derivado de los objetivos fijados en el proyecto de investigación.

El primer paso consiste en fijar el nivel de análisis. Por nivel de análisis se entiende el grado con el que el observador define su unidad de análisis en función de sus necesidades teóricas y metodológicas.

A continuación, y a modo de ejemplo, se presentan niveles de análisis y posibles alternativas, para la investigación por observación de un proceso de compra.

Tabla 1

Nivel de análisis	Alternativa
Individuo	niño, hombre, mujer, etc.
Establecimiento	gran superficie, hipermercado, supermercado, tienda tradicional, etc.
Grupo	familia, amigos, etc.
Diada	pareja, madre hijo
Proceso de compra	reflexivo, impulso
Cultura	latino, anglosajón, magrebí, etc.

La elección de una alternativa u otra dentro de un determinado nivel de análisis, tiene implicaciones de tipo conceptual.

Una vez definido el marco teórico, esto es, el conjunto de principios teóricos que guían la investigación estableciendo las unidades de análisis relevantes para cada problema de investigación, deberemos definir las categorías de observación; por éstas se entienden las definiciones operativas de clases que permiten registrar el fenómeno bajo observación.

Las categorías de observación se dividen en:

- **Categorías excluyentes.** Representan el conjunto de categorías que cumplen el requisito de no solapar su contenido, de modo que cada uno de los elementos del

fenómeno bajo observación sólo pueden ser registrados dentro de una categoría de ese conjunto.

- Categorías exhaustivas. Conjunto de categorías que cumplen el requisito de abarcar todos los elementos que componen el fenómeno bajo observación.

2.3 TIPOS DE OBSERVACIÓN

Existen diversos métodos de clasificar la investigación por observación. Los más habituales son:

2.3.1 EN FUNCIÓN DE LA INTERVENCIÓN DEL INVESTIGADOR

Podemos clasificar la observación en:

Observación natural, estructurada y experimento de campo, observación participante y auto observación

Observación natural

En este tipo de estudio, el investigador es un mero espectador de la situación y no interviene en modo alguno en el curso de los acontecimientos observados.

Se produce dentro del contexto usual en el que surgen los fenómenos de interés para el investigador.

Observación estructurada

Es un plan de recogida de datos mediante observación, llevada a cabo en el contexto natural en el que se produce el fenómeno que se quiere observar y en el que el investigador trata de establecer algún tipo de control sobre la situación.

Experimento de campo

Es un experimento realizado en una situación natural. El experimento de campo conlleva la creación de, al menos, dos situaciones diferentes de observación, de modo tal que las diferencias que se espera que aparezcan entre ambas sean atribuibles a la causa cuyo influjo se está investigando. Es preciso disponer de una teoría tentativa que explique los datos que se obtengan de la observación.

Observación participante

El observador forma parte de la propia situación bajo observación.

Auto observación

Es un tipo especial de observación en la que observador y observado son la misma persona.

Investigación Comercial

Proviene de los primeros experimentos de laboratorio de Wundt y se denominaba introspección. En esta técnica se estructuraban una serie de situaciones en el laboratorio y se pedía a los investigados que relataran sus experiencias subjetivas, constituyendo éstas los datos de la investigación.

En la actualidad se utiliza para el registro de conductas de tipo interno. Por ejemplo, podemos pedir al consumidor que anote el producto cada vez que lo use o consuma y, a la vez, que registre su grado de satisfacción o sus sentimientos, etc.

2.3.2 SEGÚN SE REALICE EN CIRCUNSTANCIAS NATURALES O ARTIFICIALES

Observación artificial

El investigador manipula o altera deliberadamente el ambiente con el objeto de crear una situación particular y observarla. La aplicación de esta técnica es interesante cuando algunos comportamientos no se presentan frecuentemente y el coste de esperar a que ocurra es prohibitivo.

2.3.3 SEGÚN LA PARTICIPACIÓN DE LA MUESTRA

Es decir, si son conocedores o no de que se les está investigando.

Observación encubierta

El observado no sabe que está siendo objeto de investigación. Se recurre a ella siempre que se considera que la persona a la que vamos a observar se comportaría de forma diferente si supiera que se la está observando.

Observación no encubierta

Se solicita la participación de la muestra. Por ejemplo: un panel de audiómetros.

2.3.4 DE ACUERDO CON SU ESTRUCTURACIÓN

Según sea estructurada o no

Observación estructurada

Conocemos de antemano los tipos de actividades y las características que identificamos y registramos.

Observación inestructurada

El investigador registra cuanto estima pertinente del hecho investigado. Es buen procedimiento como investigación exploratoria.

2.3.5 EN FUNCIÓN DE LA FORMA

Según se realice de forma directa o indirecta

Observación directa

Se contempla el comportamiento del investigado tal y como se realiza. Un ejemplo de esta técnica puede darse en una clase, al seleccionar una muestra y contar el número de veces que miran el reloj. Cuando una clase resulta aburrida, la gente tiene tendencia a mirar la hora.

Observación indirecta

Consiste en ver los resultados de un comportamiento ya realizado. También se llama estudio de trazas o residuos físicos. Ejemplos de esto serían la técnica de Dustbin Check. La auditoria de despensa (se pide permiso para examinar las casas de los participantes en busca de ciertos productos o marcas). Estudio de contenido: consiste en identificar las características específicas que existen o no en los materiales bajo estudio. Un ejemplo de este tipo de estudios sería estudiar el papel de la mujer, examinando la publicidad en revistas destinadas al público en general.

2.3.6 DE CONFORMIDAD CON EL PROCEDIMIENTO

Según la procedencia de la observación : humana o mecánica

La observación en el primer caso es realizada por personas y en el segundo, se emplean procedimientos mecánicos o electrónicos. Los medios mecánicos más utilizados son: audiómetros, cámara ocular (fotografía el movimiento de los ojos), psicogalvanómetro (máquina de la verdad), cámara de vídeo, pupilómetro (mide el diámetro de la pupila) y taquitoscopio.

2.4 ELABORACIÓN DE UN CÓDIGO DE OBSERVACIÓN

Para el desarrollo de este apartado vamos a partir del principio de que, al iniciar la investigación, no disponemos de ningún punto de referencia anterior, es decir, carecemos de fuentes secundarias.

En esta situación seguiremos la metodología de Bakeman y Gottman, los cuales en su obra *Observación de la interacción: Introducción al análisis secuencial* (Ediciones Morata, 1989), plantean las recomendaciones que resumimos seguidamente:

Nunca se debe comenzar a observar algo si no se tiene, previamente, una pregunta que responder. La pregunta deberá estar formulada de una manera clara, concisa y

Investigación Comercial

concreta (CCC). La propia formulación de la pregunta incluirá en sí misma una respuesta tentativa, ya que definirá el problema en un ámbito concreto.

Una vez formulada la pregunta, se debe elegir el nivel o niveles de análisis adecuados para tratar de encontrar una respuesta. Se trata de diseñar categorías que tengan sentido para el problema objeto de estudio.

Dedicar un tiempo previo a realizar una observación asistemática. En esta fase recogeremos la información de una forma narrativa. Este tipo de observación asistemática nos servirá como enfoque (investigación exploratoria), para la posterior realización de la observación sistemática (OS), permitiéndonos seleccionar las categorías.

Hay que utilizar categorías que estén dentro del mismo nivel de definición (molaridad y molecularidad)¹, que sean homogéneas y con el suficiente nivel de detalle para el problema en cuestión. La homogeneidad hace referencia a que los eventos incluidos en una categoría sean lo más parecidos posible. El nivel de detalle trata de evitar que descubramos, sobre la marcha, nuevos detalles que puedan ser interesantes y relevantes, que no estemos registrando.

El código debe estar compuesto por categorías exhaustivas y excluyentes entre sí. En muchas ocasiones, esta circunstancia no es posible.

Tras el registro narrativo obtenido como consecuencia de la observación asistemática, llegaremos a la propuesta del código. Este código necesitará a su vez un proceso de depuración, mediante contrastación empírica (test de prueba), antes de considerarlo como válido para el objetivo de la investigación.

Una vez definido el código con sus diversas categorías, exhaustivas y excluyentes, deberemos comenzar la observación y realizar los correspondientes registros.

2.5 MEDIDAS EN LA INVESTIGACIÓN POR OBSERVACIÓN

En las técnicas de observación se diferencian cinco tipos de medidas diferentes y que son las siguientes:

Ocurrencia, frecuencia, latencia, duración, intensidad.

¹ Los psicólogos utilizan la terminología química: molaridad y molecularidad, para referirse a niveles generales (molar) o más específicos sencillos (molecular)
Investigación Comercial

Todas ellas hacen referencia al fenómeno que estamos registrando dentro de cada categoría. La elección del tipo de medida depende tanto de la naturaleza del fenómeno bajo observación como de los intereses del propio investigador.

Veamos en qué consiste cada una de estas mediciones.

Ocurrencia: Nos informa si determinado fenómeno o suceso aparece o no durante el periodo de observación.

Frecuencia: Es el número de veces que sucede un determinado dato de observación durante el periodo de observación.

La frecuencia a su vez se divide en absoluta y relativa.

Latencia: tiempo que transcurre desde la aparición de un estímulo y la reacción ante el mismo.

Duración: tiempo durante el que se manifiesta el fenómeno objeto de la investigación.

Intensidad: fuerza con la que el fenómeno que estamos observando aparece en un momento dado de la observación.

2.6 TÉCNICAS DE MUESTREO Y OBSERVACIÓN

En la investigación por observación, al igual que en el resto de las técnicas de la investigación comercial, se nos plantea el problema de “a quién observar”. Podemos observar a todos los componentes de la población o universo objeto de estudio, circunstancia que en la mayor parte de los estudios de mercado y de Marketing es imposible o bien centrarnos en una muestra estadísticamente representativa de la población.

Para la selección de la muestra se aplican las técnicas de muestreo en sus variantes de muestreo aleatorio simple y sistemático.

Una vez elegida la muestra que se va a estudiar, el paso siguiente es definir cuándo realizar la investigación, esto es, hacerla continua durante todo el periodo de tiempo (por ejemplo, todo el día), o en determinadas fracciones. Esta selección del tiempo dedicado a observar es lo que en esta técnica se denomina *muestreo de tiempo*.

Por muestreo de tiempo entendemos la acción de extraer muestras de una población compuesta por intervalos de tiempo.

Existen dos procedimientos para realizar el muestreo de tiempo: muestreo sistemático y muestreo aleatorio.

Muestreo sistemático: en este caso el investigador, mediante un criterio racional, selecciona los periodos de observación.

Muestreo aleatorio: en este tipo de muestreo se divide el tiempo en fragmentos, eligiéndose éstos mediante un procedimiento de azar. Todos los fragmentos tienen que tener la misma probabilidad de ser elegidos, obteniéndose por tanto una muestra representativa de todos ellos.

Una vez seleccionados los periodos de tiempo durante los que se va a observar, deberemos decidir el tipo de registro que se va a realizar.

El registro de datos puede efectuarse de una forma continua o a intervalos. De forma continua implica que, durante el periodo de tiempo seleccionado, se registra en todo momento. En cambio, en el registro por intervalos se establece un procedimiento mediante el cual se intercalan intervalos de observación con intervalos de registro. Se coloca algún dispositivo avisador para que el investigador actúe observando o registrando.

En una gran cantidad de casos, el investigador trata de alcanzar un compromiso entre la situación ideal y las posibilidades de llevar adelante la investigación. Es decir, tiene que decidir dónde investigar, con el fin de garantizar que la situación de observación sea representativa de todas las situaciones no observadas a las que se pretende generalizar los resultados. Por ello se habla de muestreo de situaciones.

2.7 EVALUACIÓN DE LA OBSERVACIÓN

La observación de tipo científico se diferencia de cualquier otro tipo de observaciones en que debe ser sistemática, es decir, debe dar como resultado datos susceptibles de ser replicados por cualquier otro investigador.

Para analizar si la investigación por observación está bien realizada se barajan los conceptos de fiabilidad y validez

Fiabilidad.

Por fiabilidad se entiende el criterio a seguir para la valoración de un sistema de recogida de datos que nos informa del grado en el que dos investigadores, dada una misma situación, obtienen los mismos resultados. También es el grado en el que un investigador, dada la misma situación, obtiene los mismos resultados en dos momentos diferentes.

La fiabilidad hace por tanto referencia al hecho de que un procedimiento de recogida de datos nos lleve siempre a la obtención de la misma información, dentro de una determinada situación, independientemente de quien recoja los datos o del momento de su recogida.

Validez.

Por validez se entiende que el procedimiento de recogida de los datos tome aquellos que realmente pretende recoger.

Técnicas para el estudio de la fiabilidad

La fiabilidad implica que el procedimiento que estemos utilizando dé lugar siempre a los mismos resultados, cuando se cumplan las mismas condiciones.

En el caso de la técnica de investigación por observación, en la que no se replican los resultados, es complicado conocer si lo que falla es la fiabilidad del procedimiento o la naturaleza de los datos en sí mismos.

Para dar respuesta a esta circunstancia se recurre al concepto de “grado de acuerdo” entre dos observadores. Existen diversos procedimientos para el cálculo del grado de acuerdo; los más usuales por su sencillez son:

porcentaje de acuerdo y coeficiente Kappa de Cohen

Ambos procedimientos son de fácil aplicación en las mediciones de ocurrencia y de frecuencia.

Se utilizan coeficientes de correlación cuando lo que medimos son variables cuantitativas y en caso de mediciones, de latencia, duración e intensidad.

Porcentaje de acuerdo y coeficiente Kappa de Cohen

El porcentaje de acuerdo entre dos observadores es un índice que nos indica, en valor de tanto por ciento, las veces en que dos observadores han coincidido en sus observaciones sobre el total de observaciones realizadas sobre el mismo fenómeno.

Tiene el inconveniente de que no contempla la posibilidad de que algunos acuerdos puedan deberse al azar. Cohen, para paliar este efecto, propuso introducir un factor corrector, que es el coeficiente Kappa.

Para efectos prácticos se consideran fiables aquellos estudios que tengan al menos un 80% de acuerdo entre dos observadores.

El porcentaje de acuerdo entre dos observadores viene dado por la siguiente expresión matemática:

Investigación Comercial

$$P_0 = \frac{n^{\circ} de. acuerdos}{n^{\circ} de. acuerdos + n^{\circ} de. desacuerdos} \times 100$$

El coeficiente Kappa viene dado por la expresión siguiente:

$$K = \frac{P_0 - P_e}{1 - P_e}$$

Donde

$$P_0 = \frac{n^{\circ} de. acuerdos}{n^{\circ} de. acuerdos + n^{\circ} de. desacuerdos} \times 100$$

P_e representa la proporción de acuerdos esperados por azar, y que viene dado por la relación siguiente:

$$P_e = \sum_{i=1}^n (p_{i,1} \cdot x \cdot p_{i,2})$$

Donde n es el número de categorías, “ i ” es el número de la categoría, p_{i1} es la proporción de ocurrencia de la categoría “ i ” para el observador 1 y p_{i2} es la proporción de ocurrencia de la categoría “ i ” para el observador 2.

2.7.1 CASO PRÁCTICO

Supongamos dos observadores trabajando en un sistema de tres niveles (1, 2 y 3), durante un trabajo de 10 intervalos. Los resultados obtenidos son:

Intervalo	1	2	3	4	5	6	7	8	9	10
Observador 1	2	1	1	2	3	3	1	1	2	3
Observador 2	2	3	1	2	3	2	1	1	2	3

El porcentaje de acuerdo será

$$P_0 = \frac{n^{\circ} de. acuerdos}{n^{\circ} de. acuerdos + n^{\circ} de. desacuerdos} \times 100$$

Como el número de acuerdos es 8 y el número de desacuerdos es 2, sustituyendo obtenemos

$$P_0 = \frac{8}{8 + 2} \cdot x.100. = 80\%$$

Análisis de la Investigación Cuantitativa

luego, aparentemente, este estudio sí que sería fiable, ya que existe un grado de acuerdo del 80%. Ahora bien, si tenemos en cuenta la posibilidad de que algunos acuerdos sean al azar, aplicaremos el factor de corrección del coeficiente de Kappa para paliar el error.

El procedimiento que hay que seguir es el siguiente:

Con los datos de los observadores construimos la correspondiente matriz. El número de acuerdos entre los dos investigadores aparece en la diagonal de la matriz; fuera de esta diagonal quedan los desacuerdos.

En los márgenes de la matriz aparecen los porcentajes de ocurrencia de cada uno de los tres niveles según cada investigador.

		Observador 2			
		Nivel 1	Nivel 2	Nivel 3	
Observador 1	Nivel 1	3	0	1	4 / 10
	Nivel 2	0	3	0	3 / 10
	Nivel 3	0	1	2	3 / 10
		3 / 10	4 / 10	3 / 10	

El coeficiente de Kappa viene dado por

$$K = \frac{P_o - P_e}{1 - P_e}$$

En este caso P_o es 0'8 y

$$P_e = \left(\frac{4}{10} \times \frac{3}{10}\right) + \left(\frac{3}{10} \times \frac{4}{10}\right) + \left(\frac{3}{10} \times \frac{3}{10}\right) = 0'33$$

Sustituyendo, obtenemos que

$$K = \frac{0'8 - 0'33}{1 - 0'33} = 0'70$$

Luego $K = 70\%$, es decir, en este estudio, el porcentaje debido al azar es del 10%. Como no se alcanza el nivel del 80%, este estudio no sería fiable.

2.8 VALIDEZ

Se trata de establecer un criterio para la valoración de un sistema de registro de datos que nos informa del grado en el que el sistema consigue observar lo que pretendía.

Investigación Comercial

El grado de validez, en su nivel estadístico, se realiza a través de índices de correlación entre las diferentes medidas del mismo suceso.

La validez se suele estudiar con relación a tres aspectos diferentes:

validez de contenido, validez de constructo y validez orientada al criterio

2.8.1 VALIDEZ DE CONTENIDO

Indica el grado en el que los elementos incluidos en el código de observación son representativos de todo el fenómeno bajo observación.

2.8.2 VALIDEZ DE CONSTRUCTO

Indica en qué medida un código de observación es congruente con la teoría desde la que se elaboró.

2.8.3 VALIDEZ ORIENTADA AL CRITERIO

Indica el grado en que el código de observación es sensible a las variaciones del fenómeno bajo observación.

2.9 ANÁLISIS DE DATOS

Se utilizan los conocidos indicadores de tendencia central: moda, mediana y media. Como medidas de dispersión: rango, desviación intercuartil y desviación típica. Todos estos tipos de medida hacen referencia a una única variable pero, normalmente, en este tipo de investigación comercial, lo que interesa es buscar y establecer la relación entre las diferentes variables. Para ello se estudian los índices de correlación. La correlación nos indica la covariación. La variación simultánea de dos variables no indica necesariamente que una influya en la otra. Cuando sí que existe influencia, hablamos de la existencia de relación causal.

2.10 FUENTES DE ERROR EN LA OBSERVACIÓN

Para realizar una buena investigación por observación, lo importante es que el investigador no interfiera en el curso de los acontecimientos.

Está demostrado que simplemente la mera presencia de un observador, altera el comportamiento usual de las personas. Este tipo de sesgo se conoce como “reactividad”. Este tipo de sesgo se elimina logrando, o bien que el observador pase totalmente desapercibido, o bien utilizando medios mecánicos o electrónicos tales como: cámara

oculta, espejos unidireccionales, etc. Cuando estos procedimientos no se pueden aplicar, se espera a que el observado se acostumbre a la presencia del investigador; ello implica realizar sesiones de observación cuyos datos son desechados.

Otra fuente de error es el sesgo por parte del investigador, que es consecuencia de que todo investigador, al realizar una hipótesis de trabajo y recoger los datos de una determinada situación, tiende a sesgar su percepción en favor de sus expectativas.

El sistema de dar solución a este tipo de sesgo es que el “observador sea ciego”. Por observador ciego se entiende que el investigador que tiene que realizar el estudio no conoce la razón por la que se está observando, simplemente se le entrena en el uso del código.

2.11 VENTAJAS E INCONVENIENTES DE LA INVESTIGACIÓN POR OBSERVACIÓN

De una forma esquemática las podemos resumir en:

Ventajas

La información recogida, por regla general, es muy objetiva. Ello es debido a la eliminación de las influencias sobre el observado.

No requiere, por lo general, la colaboración de las personas investigadas.

Se obtiene información sobre la conducta del individuo, en ocasiones incluso desconocida por la misma persona objeto de investigación.

Inconvenientes

No se obtiene información acerca de las opiniones de los consumidores, ni de preferencias, motivos de compra, etc.

Por lo general su coste es elevado, requiere de personas cualificadas, así como de instrumental, también costoso.

Es un procedimiento lento.

3. EXPERIMENTACIÓN

3.1 INTRODUCCIÓN

En este tema se estudia cómo, a través de la experimentación, se puede analizar el efecto que una variable independiente produce sobre otra variable dependiente. Para ello es necesario controlar y neutralizar la influencia que otros factores puedan ejercer sobre la variable objeto de estudio; con este fin, la experimentación se traspa a universos aleatorios en los que el control es probabilístico y los resultados obtenidos se estudian a través del análisis de la varianza.

3.2 CONCEPTOS GENERALES

Uno de los objetivos de la Investigación de Marketing es el de tratar de definir las relaciones que unen al mix de Marketing de la empresa con sus resultados.

Esta información es de suma importancia en el proceso de toma de decisiones, así como en la planificación estratégica y en los mecanismos de control de la misma.

Las relaciones que se identifican entre las variables del Marketing mix de la empresa y sus resultados son de tipo causa efecto, constituyendo lo que se denomina relaciones de causalidad.

El análisis causal es el que pretende investigar las relaciones de influencia o causalidad entre las diferentes variables.

Desde un punto de vista filosófico, se puede entender como causa aquello que hace ser a algo que no es, o que venga a ser de forma distinta lo que es. Este concepto de causa implica el que se diferencie entre la causa que produce algo nuevo de la que sólo modifica lo existente.

Teniendo en cuenta que la investigación de Marketing no se ocupa de los consumidores y productos en su conjunto, sino sólo de las variables de éstos en los estudios descriptivos y de las relaciones entre las variables en los explicativos, es obvio que a la investigación de Marketing le interesa la causalidad no en el sentido que produce un nuevo ser, sino en la modificación de lo existente.

Cuando se dice que dos variables están unidas por una relación de causalidad, significa que una variable influye en la otra, en el sentido de que una modificación en la primera conduce a una variación en la segunda.

Investigación Comercial

Análisis de la Investigación Cuantitativa

Para el estudio de la causalidad se utilizan métodos correlacionales y experimentales. Seguidamente vamos a estudiar los métodos experimentales.

El método de experimentación consiste en reproducir fenómenos a voluntad del investigador. Aplicado a la Investigación de Marketing, trata de provocar la conducta del consumidor en condiciones perfectamente controladas, lo más parecidas posible a una situación real, con el objetivo de sacar consecuencias de la respuesta a un estímulo cuyo efecto queramos conocer.

Es un método de investigación que ayuda a identificar relaciones causales (causa/efecto). Mide el efecto que produce la variable independiente sobre la dependiente (por ejemplo, promoción de un producto y ventas obtenidas).

La principal dificultad de la experimentación consiste en realizar la prueba en las mismas circunstancias que en la realidad, así como en aislar los resultados obtenidos, debido a la variación producida respecto a otras variables no controladas en el experimento. Lo que hacemos es introducir modificaciones en variables de Marketing y tratamos de controlar su incidencia en el comportamiento de compra por parte de los usuarios.

La ventaja de este método es que elimina el factor distorsionador que el entrevistado provoca al suministrar información en una encuesta, ya que lo que aquí se estudia es el comportamiento del consumidor ante una determinada situación.

Ejemplo: Investigación sobre el color de los envases y su influencia sobre el precio y las ventas.

La experimentación podría realizarse en una serie de tiendas. En unas se venderían de una determinada característica y en otras, con otro tipo de característica. Al cabo de un período de tiempo se alternarían las características, y posteriormente se analizarían los resultados. Si el diferencial fuese significativo, podríamos obtener como consecuencia cuál es la característica preferida por el consumidor.

La experimentación puede ser más compleja, y entonces se recurre al denominado *mercado de prueba* (reproduce las condiciones del mercado corriente pero en una zona reducida).

La realización del método de experimentación debe ser perfectamente planificada.

Los aspectos básicos de esta planificación son:

1. Definición de los objetivos

Investigación Comercial

2. Definir la zona experimental
3. Elección al azar de las unidades experimentales
4. Período de duración de la experimentación
5. Diseño experimental
6. Recogida de información

Definición de los objetivos

Deben estar claramente establecidos, especificando la variable cuya influencia en otra se trata de comprobar.

Determinar la zona experimental

Debe reunir las condiciones de representatividad necesarias.

Duración de la experimentación

Depende de la frecuencia de compra. A modo de ejemplo:

compra cotidiana _____ mínimo un mes

compra ocasional _____ mínimo 5 meses

Elección del diseño

Los diseños son muy variados; radica en la forma de atribuir los diferentes tratamientos que se quieren probar a las unidades experimentales elegidas.

3.3 PRINCIPALES APLICACIONES

Entre las aplicaciones más utilizadas podemos reseñar las siguientes:

- Fijación de precios
- Selección de medios publicitarios y promocionales
- Elección de puntos de venta
- Determinación del tipo de envase y su tamaño
- Lanzamiento de nuevos productos

3.4 METODOLOGÍA

En todo experimento habrá que definir:

1 Factor principal: variable independiente estudiada con sus diferentes alternativas, a las que se denomina “tratamientos”.

2 Factores externos: son los factores influyentes que es conveniente aislar y controlar.

En algunos diseños experimentales se estudian de forma individual y se denominan factores bloque o rodeo.

Investigación Comercial

3 Unidades experimentales: lugares donde se realiza el experimento. Se dividen en los siguientes tipos:

- a) de laboratorio: local donde se reproducen las condiciones reales del mercado. Normalmente se suele hacer en el propio centro de investigación.
- b) natural o real: el estudio se realiza en lugares muestra del mercado real, zonas geográficas, ciudades, tiendas, etc.

4 Variable dependiente: variable de respuesta por parte del mercado. Nos permite medir los efectos de las variables estudiadas.

Ejemplo :

Una empresa de conservas vegetales desea medir el efecto de dos estrategias de promoción diferenciadas para comercio en régimen de autoservicio y para tiendas especialistas. Definir las características del experimento.

- 1 Factor principal: los dos tipos de promoción
- 2 Factor externo: situación del producto en la tienda, en la estantería, día de la semana
- 3 Unidad experimental: comercio de las características requeridas (tiendas reales)
- 4 Variable dependiente: unidades físicas de producto vendidas

3.5 INDICADORES DE VALIDEZ DE UN EXPERIMENTO

La validez de un experimento comercial viene dada por dos indicadores:

- 1. Validez interna: capacidad de aislar los efectos del factor estudiado.
- 2. Validez externa: capacidad de generalizar los resultados del experimento.

Por lo general, el incremento de validez de un experimento va asociado con un aumento en los costes económicos del estudio.

Por consiguiente, el objetivo será encontrar un diseño experimental que equilibre el coste económico con el nivel de validez.

3.6 ANÁLISIS ESTADÍSTICO DE LOS DISEÑOS EXPERIMENTALES

Las técnicas estadísticas más usuales en la experimentación comercial son:

El análisis de la varianza ANOVA.

El análisis de la covarianza ANCOVA.

Cuando se quieren medir los efectos de uno o más factores en dos o más variables dependientes se utiliza el análisis de la varianza, denominado MANOVA.

El análisis de la varianza (ANOVA) se utiliza cuando no conocemos previamente el comportamiento de la variable dependiente sin la influencia del factor principal controlado.

En esta situación, se realiza un test estadístico para medir los efectos del factor sobre la variable.

El análisis de la covarianza (ANCOVA) se utiliza cuando conocemos previamente, a través de mediciones previas o de grupos o unidades de control, el comportamiento de la variable dependiente sin la influencia del factor principal controlado. De esta manera se evita la posible influencia de factores externos no controlados. El test estadístico que se va a realizar consta de dos partes:

Se comprueba la existencia de factores externos no controlados.

Si no existen, se miden los efectos del factor principal en la variable dependiente a través del análisis de la varianza.

Los análisis de la varianza y de la covarianza se pueden realizar a través de dos métodos: el denominado Tradicional (anova) o bien mediante la Regresión

3.6.1 MÉTODO ANOVA TRADICIONAL

El proceso de este método es:

Se determinan las siguientes dispersiones:

- 1.- Dispersión total (DT): mide la suma de las dispersiones.
- 2.- Dispersión factorial (DF): mide la dispersión entre los grupos creados por las diferentes alternativas del factor o factores estudiados.

Dependiendo del tipo de experimento, pueden existir varias dispersiones factoriales, correspondientes al factor principal y a los factores de bloque.

- 3.- Dispersión residual (DR): mide la dispersión dentro de los grupos creados por las diferentes alternativas del factor o factores estudiados.

$$DT = DF + DR \quad DR = DT - DF$$

- 4.- Se calcula el cuadrado medio total (CMT): dispersión total dividida por el número de grados de libertad.

$$CMT = DT / gl \quad \text{donde } gl \text{ son los grados de libertad}$$

5.- Se calcula el cuadrado medio factorial (CMF): dispersión factorial dividida por el número de grados de libertad.

$$\text{CMF} = \text{DF} / \text{gl}$$

dependiendo del tipo de experimento, pueden existir varias varianzas factoriales, correspondiendo éstas al factor principal y a los factores bloque.

6.- Se calcula el cuadrado medio residual (CMR): dispersión residual dividida por el número de grados de libertad.

$$\text{CMR} = \text{DR} / \text{gl}$$

7.- Se realiza el test de la F: para cada factor estudiado se calcula el correspondiente valor del estadístico F.

7-1.- Se calcula el estadístico F para cada factor objeto de estudio, la ecuación correspondiente es:

$$F = \text{CMF} / \text{CMR}$$

Si el valor de F es menor que uno, es decir, $\text{CMF} < \text{CMR}$, no existe un efecto significativo del factor estudiado sobre la variable dependiente, y por tanto no es necesario realizar la comparación de F con el correspondiente valor de las tablas.

7-2.- Se determina el valor de F en las tablas estadísticas de la distribución de la F, en base a los grados de libertad del numerador y del denominador.

7-3.- Se comparan ambos valores.

La hipótesis nula H_0 es: NO EXISTE EFECTO SIGNIFICATIVO DEL FACTOR ESTUDIADO.

Entonces:

Si $F > F_t$ (tabla), no se cumple H_0 y por tanto el factor estudiado tiene una influencia significativa sobre la variable dependiente.

Si $F = F_t$ (tabla), no podemos rechazar H_0 .

3.7 TIPOS DE EXPERIMENTOS HABITUALMENTE UTILIZADOS EN INVESTIGACIÓN COMERCIAL

Los tipos de experimentos que más habitualmente se utilizan en la Investigación Comercial son:

3.7.1 EXPERIMENTACIÓN AL AZAR

En este tipo de experimento comercial sólo se controla un factor: la variable independiente estudiada.

La asignación de tratamiento a las diferentes unidades experimentales se realiza de forma aleatoria.

3.7.2 EXPERIMENTACIÓN EN BLOQUES ALEATORIOS

En este tipo de experimento comercial se controlan dos factores:

3.7.3 EXPERIMENTO DE CUADRADO LATINO

En este tipo de experimentación comercial se controlan tres factores:

- 1 La variable independiente o factor principal
- 2 Dos factores de control o rodeo que se denominan “factores bloque”

El diseño en cuadrado latino exige utilizar el mismo número de alternativas en los tres factores controlados.

Debemos plantear este tipo de estudio cuando se estima que existen otros dos factores influyentes en el fenómeno estudiado, aparte del factor principal.

Se debe diseñar un número de unidades experimentales suficiente para probar todas las combinaciones posibles entre los tres factores sometidos a control.

3.7.4 EXPERIMENTO CON INTERCAMBIO

Este procedimiento consiste básicamente en la aplicación alternativa y sucesiva de los diferentes tratamientos a las unidades experimentales. El orden de aplicación de los diversos tratamientos sobre las unidades experimentales debe ser al azar, con la condición de que haya el mismo número de unidades experimentales que reciba primero un tratamiento y después los otros.

Este tipo de experimento combina las características de los bloques aleatorios y los de los cuadrados latinos pequeños.

3.7.5 EXPERIMENTO FACTORIAL

En los experimentos comerciales de tipo factorial se controlan varios factores principales, midiendo sus efectos individuales y los conjuntos sobre la variable dependiente.

Esta es una situación muy habitual en el área de Marketing, donde la aplicación del Marketing mix produce en el mercado unos resultados diferentes de los que se obtendrían por la suma de los efectos aislados de cada factor del mix de Marketing.

La técnica estadística que se utiliza se denomina ANOVA de vía múltiple.

3.8 LIMITACIONES DE LA EXPERIMENTACIÓN COMERCIAL

Las aplicaciones de la experimentación comercial presentan las siguientes limitaciones:

Aplicación a corto plazo. La mayor validez de la experimentación comercial es a corto plazo, ya que en este período las condiciones y circunstancias bajo las que se realiza la experimentación sufren poca variación.

No es una técnica adecuada para estudiar productos de baja frecuencia de compra. En la medida que los productos son adquiridos por los consumidores con gran asiduidad, la experimentación puede realizarse en periodos muy cortos, con lo que los resultados que se obtienen son más valiosos.

Dificultad de aislar el mercado de prueba. Cuando se realiza una experimentación en una zona, en diversas tiendas, etc., resulta difícil evitar que se produzcan distorsiones en la zona de prueba, como consecuencia de las compras efectuadas por los consumidores.

Destrucción por actuaciones de la competencia. Esta técnica puede hacer que la competencia conozca la investigación que se está realizando.

Coste. La experimentación, por lo general, tiene un coste elevado.

3.9 CASO PRÁCTICO: EXPERIMENTO DE BLOQUE ALEATORIO

Recordemos que en este tipo de experimento comercial se controlan dos factores:

1 La variable independiente o factor principal

2 Un factor de control que se denomina “factor bloque”, también llamado “de rodeo”

CASO: Una empresa de refrescos va a lanzar al mercado un nuevo producto; para ello realiza una prueba con tres envases diferentes:

P1 envase de 2l., P2 envase de 1l., P3 envase de 0’5l.

Además, la empresa controla otro factor influyente, el tipo de establecimiento donde se expenden los refrescos; para ello, definen el siguiente factor bloque:

B1 grandes superficies, B2 supermercados, B3 tienda tradicional y B4 autoservicio.

Cada envase se prueba en los cuatro tipos de tienda, durante un mes. Se obtienen los resultados siguientes en miles de unidades de producto:

Investigación Comercial

Análisis de la Investigación Cuantitativa

Tabla de resultados:

	B1	B2	B3	B4
P1	3	4	3	2
P2	7	8	7	6
P3	8	12	8	4

SOLUCIÓN:

Factor principal: tratamientos P1, P2, P3, luego $k=3$

Factor bloque: las alternativas B1, B2, B3, B4, luego $R = 4$

Unidades experimentales $4 \times 3 = 12$

Variable dependiente: unidades vendidas

Siendo

n el número de mediciones (12)

x_{ij} las unidades vendidas en los diferentes establecimientos

m_j la media de ventas por tratamiento

m_i la media de ventas por cada alternativa de bloque

m la media total

Cálculos

	B1	B2	B3	B4	S	m_j
P1	3	4	3	2	12	3
P2	7	8	7	6	28	7
P3	8	12	8	4	32	8
S	18	24	18	12		
m_i	6	8	6	4		

Luego $m = 6$

Dispersión total $DT = 92$

Dispersión factorial principal $DF = 56$

Dispersión bloque

$$DB = \sum k(m_i - m)^2$$

$$DB = 3(6 - 6)^2 + 3(8 - 6)^2 + 3(6 - 6)^2 + 3(4 - 6)^2 = 24$$

Dispersión residual

Investigación Comercial

Análisis de la Investigación Cuantitativa

$DR = DT - DF - DB$ Sustituyendo, $DR = 12$

Cuadrado medio factorial $CMF = 28$

Cuadrado medio bloque $CMB = 8$

Cuadrado medio residual $CMR = 2$

Test de la F

1 Factor principal

$$F = \frac{CMF}{CMR}$$

Luego $F = 14$, Como el valor en tablas para el 95% y gl 2 y 6 es 5'14

Podemos decir que existe un efecto significativo de los tratamientos estudiados para un nivel de confianza del 95%.

2 Factor bloque

$$F = \frac{CMB}{CMR}$$

Luego $F = 4$, El valor correspondiente en tablas para el 95% y gl 3 y 6 es $F = 4'76$

Como $4 < 4'76$ podemos decir que:

NO existe un efecto significativo del factor bloque para el nivel de confianza del 95%.

La correspondiente salida de SPSS es:

Análisis de la Investigación Cuantitativa

3.9.1 EXPERIMENTO DE BLOQUE ALEATORIO (SPSS)

Resumen del procesamiento de los casos

Casos					
Incluidos		Excluidos		Total	
N	Porcentaje	N	Porcentaje	N	Porcentaje
12	100,0%	0	,0%	12	100,0%

a. Ventas (miles de unidades) por Tipo de promoción, Tipo de tienda

Medias de las casillas^{b,c}

Envase Tipo de tienda		Ventas (miles de unidades)	
		Media	N
2 litros	Total	3,0000	4
1 litro	Total	7,0000	4
1/2 litro	Total	8,0000	4
Total	Gran superficie	6,0000	3
	Supermercado	8,0000	3
	Tienda tradicional	6,0000	3
	Autoservicio	4,0000	3
	Total	6,0000 ^a	12

a. Media global

b. Ventas (miles de unidades) por Envase, Tipo de tienda

c. No se han calculado las medias de orden 2 o superior debido al límite en el orden máximo de interacción.

ANOVA^a

Ventas (miles de unidades)		Método jerárquico				
		Suma de cuadrados	gl	Media cuadrática	F	Sig
Efectos principales	(Combinadas)	80,000	5	16,000	8,000	,012
	Envase	56,000	2	28,000	14,000	,005
	Tipo de tienda	24,000	3	8,000	4,000	,070
Modelo		80,000	5	16,000	8,000	,012
Residual		12,000	6	2,000		
Total		92,000	11	8,364		

a. Ventas (miles de unidades) por Tipo de promoción, Tipo de tienda

Análisis de la Investigación Cuantitativa

3.10 TABLA ESTADÍSTICA: DISTRIBUCIÓN DE LA F

NIVEL DE CONFIANZA 95%

N	m				
	1	2	3	4	5
1	161´4	199´5	215´7	224´6	230´2
2	18´51	19	19´16	19´25	19´30
3	10´13	9´55	9´28	9´12	9´01
4	7´71	6´94	6´59	6´39	6´26
5	6´61	5´79	5´41	5´19	5´05
6	5´99	5´14	4´76	4,53	4´39
7	5´59	4´74	4´35	4´12	3´97
8	5´32	4´46	4´07	3´84	3´69
9	5´12	4´26	3´86	3´63	3´48
10	4´96	4´10	3´71	3´48	3´33
11	4´84	3´98	3´59	3´36	3´20
12	4´75	3´89	3´49	3´26	3´11
13	4´67	3´81	3´41	3´18	3´03
14	4´6	3´74	3´34	3´11	2´96
15	4´54	3´68	3´29	3´06	2´90

Siendo **m** los grados de libertad del numerador y **n** los grados de libertad del denominador.

3.11 MÉTODOS DE SIMULACIÓN

3.12 CONCEPTO

El método de simulación consiste en la creación de una analogía o similitud de un fenómeno auténtico. Es una representación incompleta de la realidad que trata de duplicar el fenómeno, sin alcanzar la auténtica realidad.

Lo podemos definir como una representación incompleta del mercado y del mix de Marketing de la empresa. Se realiza mediante aplicaciones informáticas. Se utiliza fundamentalmente para obtener ideas sobre la dinámica del mercado, manipulando las

variables independientes (mix de Marketing y entorno), observando su influencia sobre las dependientes.

Un estudio de simulación necesita de entradas de información relacionadas con las características del caso a estudiar y las relaciones que están presentes. El investigador debe conceptualizar y documentar los componentes estructurales del sistema, estableciendo probabilidades para representar el comportamiento de los componentes. Los componentes, también llamados unidades, representan partes del mercado, por ejemplo: hogares, compradores, detallistas... Las variables en el sistema establecen la forma en que se comportan las unidades, estas variables están, por lo general, ligadas al Marketing mix, por ejemplo: variaciones en el precio, comunicación, producto..., así como factores del entorno, por ejemplo: competencia, demanda, coyuntura económica.

Las probabilidades se establecen en las unidades en función de cómo responden a las diferentes variables. El objetivo es que las unidades de simulación imiten el comportamiento del sistema que representan. Las conclusiones de la simulación se expresan en resultados numéricos, por ejemplo: ventas, rentabilidad, participación en el mercado....

Para realizar y desarrollar el método de simulación se precisa de personas especialistas. Su tarea consiste en reducir el fenómeno complejo que se está estudiando, a proporciones manejables, reproduciendo las interacciones que ocurren en el mercado real. Una vez construida la simulación se pone a prueba para determinar si los resultados son válidos con la realidad del mercado que estamos simulando.

Según Robert C. William T. Nevel, en su obra *Simulation in Bussiness and Economics*, “Un modelo válido de simulación debe comportarse en forma similar al fenómeno que representa. Este es un criterio de validez necesario, pero que por sí solo puede no ser suficiente para que permita confiar en sus habilidades predictivas”.

Un buen modelo de simulación debe cumplir aspectos tales como:

No ser demasiado sencillo, ni excesivamente complejo

El usuario lo debe comprender y manipular con facilidad

Será representativo del mercado sobre el que se efectúa el estudio

Deberá tener un grado suficiente de complejidad para que se aproxime al máximo a la realidad del mercado objeto de estudio.

3.13 CLASIFICACIÓN

La clasificación más usual es la basada en el propósito que cumple el modelo, reconociéndose los siguientes:

Descriptivo: es el que describe el sistema de Marketing que se está estudiando. Son fáciles de construir, su aplicación es muy limitada ya que solo tiene utilidad para la situación que representa. Generalmente, no se pueden utilizar para reproducir situaciones causales entre varias variables, sirviendo de base para el desarrollo de los otros modelos.

Predictivo: se diseña para hacer predicciones de Marketing cuando variamos las variables del sistema. Su limitación es el no permitir la manipulación del sistema para evaluar nuevos campos de acción.

Prescriptivo: permite al usuario experimentar cambios en el sistema.

3.14 VENTAJAS Y LIMITACIONES

El método de simulación exige que se conceptualice, desarrolle y manipule un modelo, obteniéndose resultados expresados en datos numéricos. En cambio, los datos obtenidos por encuestas, experimentación e incluso por trabajo de gabinete, son el resultado directo de la situación objeto de estudio.

Como ventajas podemos citar las siguientes:

Menor coste. Es más económico que otros sistemas de recogida de información. Menor tiempo para recoger y analizar la información. Permite la evaluación de diferentes estrategias de Marketing.

Es una buena herramienta de entrenamiento para directivos, tanto del área de Marketing como del resto de departamentos de la empresa. Sirve para hacer análisis de sensibilidad, esto es, determinar la sensibilidad de una opción estratégica a desviaciones de los supuestos iniciales.

Como limitaciones podemos reseñar las siguientes:

Su poca o nula viabilidad cuando la empresa no tiene antecedentes ni experiencia acerca del fenómeno que se estudia.

Necesidad de personal especialista. Dificultad de desarrollar el modelo. Continua puesta al día del modelo.

4. LA ENCUESTA ESTRUCTURADA

4.1 LA ENCUESTA. TIPOS. DATOS QUE APORTA. VENTAJAS E INCONVENIENTES

La encuesta constituye un procedimiento sistemático de recolección de datos facilitados por los entrevistados a través de cuestionarios.

Para que las encuestas tengan éxito, deben basarse en la correcta realización de un documento escrito que recoja las preguntas que se realizarán, esto es, el cuestionario.

El cuestionario debe estar redactado en un lenguaje claro y directo, evitando las preguntas ambiguas o contradictorias entre sí, que puedan ocasionar errores en las respuestas.

La encuesta estructurada recoge información sobre uno o varios temas. La información obtenida corresponde, generalmente, a una muestra de la población investigada (universo). Para que los datos recogidos a través de una encuesta puedan inferirse a toda la población o universo, se deberá obtener la información mediante técnicas de muestreo.

Las encuestas permiten obtener información sobre características socioeconómicas, opiniones, actitudes y motivaciones del público objetivo.

La encuesta puede ser diseñada específicamente para el estudio que se va a realizar, denominándose "*ad hoc*" o estándar, esto es, encuestas cuyo diseño ha sido previamente establecido y el tipo de información es uniforme para todos los suscriptores.

Para definir la naturaleza de una encuesta, podemos utilizar el criterio de la finalidad de la encuesta en relación con el nivel general de conocimientos del investigador sobre el tema que va a tratar.

Según el criterio de finalidad, las encuestas se clasifican en:

encuestas exploratorias, encuestas descriptivas y encuestas explicativas.

Veamos brevemente en qué consiste cada una de ellas.

Encuestas exploratorias

El principal objetivo es la identificación de los problemas y la reformulación más precisa de los mismos. Normalmente sirven de fase preparatoria para los otros dos tipos de encuestas.

Investigación Comercial

Análisis de la Investigación Cuantitativa

Las encuestas exploratorias requieren con frecuencia el empleo de instrumentos de investigación poco perfeccionados, presentando algunas ventajas:

Flexibilidad. Se debe, en gran medida, al reducido número de encuestados y al uso de preguntas abiertas, cuando sea necesario.

Bien realizada. Permite definir las variables explicativas.

Sus límites son que no se puede realizar análisis estadístico ni extrapolar sus resultados.

La encuesta exploratoria permite entender cómo se plantea el problema y formular hipótesis respecto a su funcionamiento.

Encuestas descriptivas

Es el tipo de encuesta más practicado en la actualidad; su objetivo principal consiste en describir las características de una determinada situación mediante el análisis de diferentes variables y obtener apreciaciones acerca del comportamiento que se trata de prever, describiendo el grado de asociación entre dichas variables.

El investigador necesita un conocimiento previo del tema tratado; hay que definir con precisión las fuentes de información, los medios de recogida y el tratamiento de las informaciones.

Para realizar una encuesta descriptiva se necesita un cuestionario estructurado y una muestra significativa del universo objeto del estudio.

Su interés se debe a que se puede realizar un análisis estadístico de los resultados, lo que permite una extrapolación al universo dentro de unos límites de fiabilidad conocidos.

Una de las principales razones de su utilización es que permite la validación de los resultados a un nivel más general; sin que por ello se puedan establecer relaciones de causa efecto entre diversas variables.

Encuestas explicativas o causales

Su objetivo principal es el de permitir prever determinados fenómenos estudiados, conociendo la naturaleza de la relación entre las causas y los efectos previstos. A lo largo de una encuesta explicativa, las dos variables discriminativas (causa y efecto) de un problema estudiado, se identifican con precisión.

Es la investigación más difícil de realizar. Las relaciones que existen entre las causas y los efectos pueden ser de dos naturalezas:

Determinista: se trata de un hecho que a la vez es necesario y suficiente para la existencia ulterior de otro suceso.

Probabilista: se trata de un acontecimiento necesario pero no suficiente para la existencia de otro suceso.

Una encuesta explicativa permite precisar y controlar los efectos frente a una causa conocida. El principal inconveniente consiste en no poder demostrar una relación exacta de causalidad entre la causa y el efecto, sino simplemente la existencia de una relación entre ambos.

Existen numerosas formas de recoger información por el método de encuestas. Cada caso requerirá la utilización de un tipo específico de encuesta. Normalmente, influyen dos criterios en la elección de las fuentes de información, la naturaleza de la encuesta y los recursos disponibles (financieros, humanos, tiempo).

4.1.1 TIPOS DE ENCUESTA ESTRUCTURADA

Hay diversos sistemas de clasificar las encuestas. Vamos a reseñar los más corrientes:

A) Teniendo en cuenta el medio empleado para obtener la información, que son las que más difusión tienen, se clasifican en los siguientes tipos básicos:

encuestas personales *vis a vis*, encuestas telefónicas, encuestas por correo y mixtas

B) Según el número de temas investigados:

Un tema: encuesta *ad hoc* (se realiza de forma esporádica, generalmente por encargo, para obtener información sobre un determinado asunto)

Varios temas: *ómnibus* (se recoge información sobre varios asuntos o temas a la vez)

C) Según el contratante

Para un determinado estudio. Se trata entonces de una encuesta *ad hoc*

Diseño por suscripción (se recoge periódicamente información de interés para varios suscriptores). Estas se clasifican en:

paneles de detallistas y paneles de consumidores

Ómnibus

D) Según la muestra:

De muestra variable. Se selecciona la muestra para cada estudio. Serían de este tipo las siguientes: *ómnibus*, encuesta sectorial de bienes de consumo duradero y encuestas *ad hoc*.

De muestra fija. La muestra permanece a lo largo de los diversos períodos de recogida de la información. En este tipo tendríamos las denominadas panel, entre las que reseñamos: panel de detallistas, panel de consumidores, panel de audiencias.

E) De acuerdo con criterios temporales:

Métodos instantáneos: son los que obtienen la información en un momento determinado de tiempo: postal, telefónica, personal, observación, experimentación. También llamados estudios transversales.

Métodos permanentes: obtienen información de una forma periódica. El sistema más utilizado es el panel. También llamados estudios longitudinales.

F) En función de la colaboración de la muestra:

Con la colaboración de la muestra. Se utiliza de forma explícita un cuestionario. A este tipo corresponden la postal, telefónica, personal, mixtas.

Sin la colaboración de los componentes de la muestra. En este caso no sólo no se pide colaboración, sino que se busca que no se den cuenta de que están siendo estudiados.

Los métodos más utilizados son: observación y experimentación.

G) En función del componente del mercado de donde se obtenga la información:

De demanda. Por ejemplo: panel de consumidores

De oferta. Por ejemplo: panel de detallistas

Como se puede observar, hay muchos sistemas de clasificación de los diversos tipos de encuestas.

4.1.2 INFORMACIÓN QUE APORTAN LAS ENCUESTAS

Los principales datos que nos permiten obtener las encuestas son:

Características del encuestado

socioeconómicas, edad, sexo, profesión, hábitat, etc.

Actitudes de los consumidores, en sus tres componentes:

afectivo: ¿qué le sugiere?

cognoscitivo: ¿conoce la marca?

comportamiento: ¿compraría?

Otras informaciones de interés para el estudio.

4.1.3 VENTAJAS E INCONVENIENTES

Sus principales ventajas son:

Investigación Comercial

Versatilidad: puede recopilar una amplia gama de información

Rapidez

Como desventajas podemos enumerar:

Resistencia del encuestado a facilitar la información pedida

Sesgo en la emisión de las respuestas, producido por el propio encuestado o influido por el entrevistador

Dificultad de recordar por parte del encuestado

4.2 CARACTERÍSTICAS DE LAS ENCUESTAS ESTRUCTURADAS

Decíamos que la encuesta es una técnica o sistema de recogida de información sobre uno o varios asuntos. La información obtenida responde a una muestra de la población objeto del estudio, aunque a veces también puede obtenerse de la propia población o universo, cuando éste es pequeño.

Las encuestas cuantitativas son un método de obtención de datos, mediante entrevista individual, en la que el entrevistado proporciona información de forma voluntaria y consciente como respuesta a una serie de preguntas planteadas en el cuestionario.

Las características básicas de la encuesta estructurada son:

- Definición del objeto de estudio
- Definición de la población que se va a estudiar (universo)
- Elaboración del cuestionario para recoger la información
- Definición del sistema de muestreo
- Definición del tamaño de la muestra en función de niveles de confianza y margen de error
- Fijación del plan de tabulación que permita el análisis de los datos obtenidos

4.3 TIPOS DE ENCUESTA ESTRUCTURADA SEGÚN EL PROCEDIMIENTO EMPLEADO PARA OBTENER LA INFORMACIÓN

4.4 ENCUESTA PERSONAL VIS A VIS

Es el procedimiento más utilizado. Se basa en entrevistas realizadas por encuestadores, debidamente entrenados, que acuden a los hogares, centros de trabajo o emplazamientos concretos, con la finalidad de realizar una encuesta. Es interactiva y debe resultar agradable tanto para el entrevistador como para el entrevistado. Para ello, el encuestador

tratará de crear un ambiente distendido, induciendo al interrogado a contestar las cuestiones.

Cuando alguien piensa en una encuesta a través de la entrevista personal, en su pantalla mental se forma un cuadro en el que ve a una persona, formulario en mano, haciendo preguntas a otra en el umbral de una puerta; si bien esta imagen es en parte cierta, cada vez su ejecución resulta más difícil, y esto es debido a varias razones, entre las cuales citaremos:

Incorporación de la mujer al mundo laboral, lo que origina unas horas de contacto fuera de los horarios laborales.

Seguridad: muchas personas son reacias a que entren extraños en sus casas, aunque vayan debidamente acreditados como entrevistadores.

Horarios de contacto con el público objetivo de la muestra, en ocasiones nocturnos o en fines de semana.

No obstante, las entrevistas personales son las que más se acercan al enfoque universal de la investigación. En teoría, y prescindiendo de aspectos económicos, pueden utilizarse en todos los estudios.

Debido a su carácter interactivo, es la técnica adecuada para tratar temas que por su complejidad precisan aclaraciones, ya que permite al entrevistador usar gráficos, escalas, fotografías, etc., es decir, cualquier elemento adicional que sea necesario. Por otra parte, permite determinar el contexto social del entrevistado.

Entre las mayores *ventajas* de este método, citaremos:

Flexibilidad y versatilidad

Es la ventaja clave de esta técnica, ya que permiten muchas opciones. Podemos obtener cuestionarios de mayor longitud y de cualquier formato. Es el método más adecuado para obtener información detallada sobre actitudes y opiniones. Se puede hacer cualquier tipo de pregunta, ya que los entrevistadores pueden hacer un trabajo de profundización y clarificación de las preguntas, y muy especialmente de las abiertas.

Demostraciones

Podemos mostrar y entregar cosas a los entrevistados. Podemos pedir su opinión sobre un anuncio, entregarle paquetes de muestra, hacer pruebas, etc.

Observación

Permite observar al entrevistado, así como su entorno.

Investigación Comercial

Rapidez

Suele ser más rápido que los estudios realizados por correo.

Se puede realizar en varios puntos geográficos a la vez.

Muestra

Es de gran utilidad en la realización de estudios en las grandes ciudades, cuando la unidad de la muestra es el hogar (método de rutas aleatorias).

En contraposición con las ventajas, las entrevistas personales presentan *desventajas* que limitan su aplicación. La más importante es su elevado coste, debido fundamentalmente a los gastos del entrevistador (viajes, dietas, repeticiones de visitas, barridos, etc.), al tiempo de ejecución, a la necesidad de una buena planificación y al amplio trabajo administrativo.

Por otra parte, no es el mejor método para tratar temas personales, íntimos o de gran sensibilidad.

Y como clara desventaja, podemos citar la inseguridad ciudadana. Debido a los horarios de trabajo del público objeto de la muestra, nos encontramos haciendo el trabajo en horarios nocturnos, encontrándolo reacio a abrir la puerta. De la misma forma, resulta difícil reclutar entrevistadores dispuestos a trabajar en determinados barrios y, especialmente, en horario de tarde-noche.

Sin embargo, para algunos tipos de investigación es necesaria la entrevista personal. La problemática indicada con anterioridad ha originado nuevas formas de contacto personal.

4.4.1 ENTREVISTAS POR DETENCIÓN

Su base es que es mucho más eficaz dejar que los encuestados vengan al entrevistador en un centro comercial, almacén, cuencas comerciales, etc., que enviar a los encuestadores a los hogares. Lo que se hace es enviar a los encuestadores a localizar a los encuestados en lugares donde sean accesibles y, una vez localizados, realizar la encuesta.

La limitación de este tipo de entrevistas es que no se genera una muestra tan representativa de la población como la obtenida en un estudio puerta a puerta. También está el problema de la duración del cuestionario, pues resulta difícil que una persona nos atienda durante mucho tiempo.

En conjunto, este tipo de estudios suele reflejar aspectos bastante cercanos a la población total si seleccionamos con rigor el punto de detención. Asimismo, nos proporciona unos costes inferiores.

Este tipo de investigación es de rápido crecimiento y cada día es más común.

4.4.2 ESTUDIOS CON INVITADOS A UNA LOCALIZACIÓN CENTRAL

En esta técnica de entrevista personal, se selecciona a los encuestados por teléfono, invitando a los que reúnan los requisitos adecuados a que concurran a un determinado sitio y hora con el fin de realizar una encuesta.

Este método es el adecuado cuando la proporción de encuestados que cumplan los requisitos es baja (se denomina *incidencia*).

Este sistema es adecuado si la entrevista es larga, supera los veinte minutos, ya que muchas personas encuentran inconveniente en dedicar en aquel momento preciso el tiempo necesario para la cumplimentación del cuestionario. En este caso podemos hacer cita previa y preselección concertando la entrevista telefónicamente.

4.5 ENCUESTA TELEFÓNICA

Conforme ha disminuido el uso de la entrevista personal, y al mismo tiempo ha crecido el número de aparatos de teléfono instalados, podemos decir que ha ido aumentando la importancia de esta técnica.

Este es el método que más está creciendo hoy en día, debido a que reduce los costes así como los riesgos derivados de la inseguridad ciudadana. Por otra parte, prácticamente la totalidad de los hogares disponen de teléfono (en España, aproximadamente un 80%).

Este sistema, de amplia difusión en EE.UU., también se está utilizando con gran amplitud en Francia, Alemania, Reino Unido, España e Italia. Esta utilización de la entrevista telefónica tiene importancia cuando existe una gran correlación entre el producto objeto de estudio y el uso o tenencia de línea telefónica. También está divulgado su uso para encuestas continuadas, cuya importancia subyace más en datos relativos que en absolutos.

El sistema de encuesta telefónica es relativamente barato y rápido. Nos permite alcanzar adecuadamente una muestra grande y dispersa. Precisa de un sistema administrativo simple y reducido. El universo está siempre disponible, pues es la Guía Telefónica; el control es muy sencillo, ya que los supervisores pueden seguir las entrevistas desde un Investigación Comercial

puesto de escucha, y por otro lado el sesgo del entrevistador tiende a ser menor que en la entrevista personal.

El sistema más desarrollado en encuestas telefónicas se denomina CATI, que corresponde a las siglas de **Computer Aided Telephone Interviewing**, que en español viene a significar Encuestas Telefónicas Asistidas por Ordenador. Mediante este sistema, el entrevistador, utilizando el teléfono, lee el cuestionario en una pantalla de ordenador, registrando las respuestas que recibe, incorporándolas a la correspondiente base de datos. Este sistema permite, en cualquier momento, obtener la salida de los resultados, lo que nos posibilita controlar continuamente el desarrollo del estudio.

El programa del sistema puede llevar a cabo una selección aleatoria de los números de teléfono.

Los requisitos físicos de los entrevistadores no son los mismos que para los de calle, siendo lo importante la dicción, el tono de voz y la forma de decir, pudiéndose emplear a personas con discapacidades físicas. En todos los casos es conveniente la realización de un diseño ergonómico del puesto de trabajo.

Las principales ventajas de la encuesta telefónica son:

Muestra: nos permite trabajar con muestras amplias y geográficamente dispersas de una manera fácil.

Las segundas llamadas pueden hacerse con mayor facilidad; ésta es una limitación importante en la entrevista personal puerta a puerta, ya que es costoso realizar un segundo contacto. En la entrevista telefónica es más fácil encontrar a la persona que no está disponible en ese momento y concertar la entrevista para otra hora; esto nos permite aplicar buenos procedimientos de muestreo a base de llamadas subsiguientes.

Supervisión: al realizarse desde una instalación telefónica central, el supervisor puede controlar el trabajo de los diferentes entrevistadores para, de esta manera, asegurarse de que el cuestionario se está utilizando de la manera correcta. Podemos hablar de control inmediato.

Flexibilidad: aunque no con la misma amplitud que la entrevista personal cara a cara, la investigación telefónica puede utilizar cuestionarios complejos, debido a que el encuestador puede controlar el interrogatorio.

Acceso a elementos de la muestra difíciles de contactar. El uso del teléfono nos permite contactar con elementos ocupados y con los alejados, con mayor facilidad que otros tipos de entrevista.

Rapidez: al utilizar un sistema centralizado, se consigue un mayor número de entrevistas por jornada, ya que no existen tiempos muertos en desplazamientos.

Permite la prueba inicial del cuestionario con gran rapidez, y hace posible la puesta en marcha en menos plazo del sondeo definitivo.

Sin embargo, la entrevista telefónica tiene una serie de *limitaciones*:

Tasa de respuesta: puede convertirse en un problema cuando la gente se da cuenta de que el tema de estudio no es importante para él, colgando el teléfono, cortando de esta manera la entrevista, situación ésta mucho más difícil en una entrevista personal *vis a vis*.

Duración de la entrevista: debe ser breve. La duración máxima debe estar entre 5 y 15 minutos. Si no se puede obtener la información en este período de tiempo, la entrevista telefónica no es la técnica adecuada.

Demostraciones: por descontado, hasta que no se masifique la utilización del vídeo teléfono, no se puede pensar en el empleo de elementos gráficos de apoyo, siendo absolutamente imposible mostrar algo al entrevistado.

Limitación en las preguntas: las preguntas deben ser sencillas y los cuestionarios cortos, siendo muy difícil manejar escalas extensas por teléfono.

La entrevista telefónica es un buen sistema para recoger opiniones, actitudes o hechos de una muestra grande y dispersa, también para estudios de seguimiento, para realizar segundas entrevistas con personas previamente contactadas de forma personal.

La entrevista telefónica es cada vez más usada para la Investigación de Marketing.

4.6 ENCUESTA POR CORREO

Esta técnica consiste en que el investigador envía por correo el cuestionario junto con una carta y un sobre franqueado, para que el encuestado envíe la respuesta.

En ocasiones, los cuestionarios se envían por correo y se recogen personalmente.

En general parece una alternativa atractiva. Es fácil de elaborar y por regla general es el sistema más barato.

Análisis de la Investigación Cuantitativa

Es un procedimiento muy conveniente para recoger opiniones de una muestra minoritaria.

En esta técnica no hay interacción personal y, por lo tanto, las preguntas deben ser fácilmente comprendidas, ya que no hay posibilidad de hacer aclaración alguna.

Las principales *ventajas* de este tipo de investigación son:

Costo: principal atracción de este método, ya que resulta mucho más económico que la entrevista telefónica y, por supuesto, que la entrevista personal puerta a puerta, así como las otras modalidades.

Eficiencia con muestras muy grandes: Conforme aumenta el tamaño de la muestra, se vuelve más eficiente; la diferencia entre enviar 2000 ó 5000 cuestionarios es el impreso y el franqueo.

Acceso a personas difíciles de localizar: facilita la comunicación con encuestados dispersos, prácticamente con el mismo costo que a personas geográficamente concentradas.

Sesgo: al no existir encuestador presente, no puede producir ninguna influencia.

Tampoco se le puede pedir al entrevistado que profundice o clarifique respuestas.

Demostraciones: por correo sí que se pueden enviar muestras, así como anexos, dibujos, etc., aunque no se puedan realizar grandes demostraciones.

Las *limitaciones* de la entrevista por correo son mucho más numerosas que las ventajas; para la mayor parte de los estudios, estas desventajas hacen que la utilización del "referendum postal" no sea la técnica apropiada.

Los principales inconvenientes son:

Baja tasa de respuesta: para un estudio por correo enviado a una serie de personas relacionadas al azar, la respuesta no supera más del 2 al 5%.

Esta falta de respuesta constituye una grave fuente de errores.

La apatía y la falta de interés son las dos razones más importantes de esta no respuesta.

El problema del desinterés puede ser atenuado con buenos resultados, enviando un segundo cuestionario a aquellos que no contestaron al primero, e incluso realizando varios envíos, pudiendo por estos sistemas alcanzar cotas de respuesta del 60 al 80%. Lo normal es que la no respuesta sea del 70 al 80% para productos de consumo y del 90 al 95% para productos industriales (datos USA. En España no suele pasar del 10%).

Otro factor que ayuda a vencer la apatía y la falta de interés es, por ejemplo, dar pequeños obsequios (participaciones de lotería, etc.) a cambio de los cuestionarios recibidos.

En cualquier caso, el índice de respuesta varía de unos países a otros, e incluso dentro del mismo país, varía de unos territorios a otros.

Sesgo provocado por la no respuesta

Aparte de la pequeña proporción de personas que responden a la encuesta, nos encontramos con que, a menudo, quienes la devuelven diligenciada no son los encuestados típicos de la muestra total.

Tampoco se puede estar seguro de quién fue la persona que realmente cumplimentó el cuestionario.

Ejemplo: puede ocurrir que nos contesten las personas a las que les gustó el producto o bien aquellos que quieren manifestar su queja.

Como anécdota podemos citar la siguiente: en 1896, Harlow Gale, de la Universidad de Minnesota, envió a 200 publicistas de Twin Cities un cuestionario postal para tratar de averiguar sus opiniones de la publicidad. No pudo realizar el estudio, ya que después de todos los intentos sólo alcanzó un 10% de respuesta.

En resumen, podemos decir que la baja tasa de respuesta y las serias dudas sobre la representatividad de los encuestados, hacen que la técnica de entrevista por correo no sea la adecuada para un gran número de estudios.

Control: al no conocer quién contesta el cuestionario, es prácticamente imposible establecer un control adecuado.

Limitaciones: las limitaciones están basadas en que el número de preguntas ha de ser pequeño, y éstas deben estar formuladas de una forma sencilla y clara. En caso contrario, los posibles encuestados desearán el cuestionario.

Otros inconvenientes son la lentitud con la que se obtienen las respuestas, así como la imposibilidad de hacer pruebas previas o estudios piloto.

Hay que concentrar los esfuerzos de esta técnica en conseguir la mayor respuesta posible.

Por tanto, el primer paso es diseñar el envío que se ha de realizar a los encuestadores.

Las principales recomendaciones que hay que tener en cuenta son:

- Calidad de los impresos

Investigación Comercial

- Tipo de correo

Si se envía un regalo, éste debe ir junto con el cuestionario, de manera que se valide con el recibo del cuestionario

- Carta de presentación, totalmente personalizada

Todo el envío debe tener aspecto personalizado y profesional; nunca debe dar la sensación de un envío al por mayor

- Incluir un sobre con el franqueo pagado
- El cuestionario tiene que parecer fácil de contestar y de enviar

Imprimir el cuestionario en los dos lados de una página doblada de 43 x 28 cm. (parece más pequeño que cuatro páginas a una cara)

Hemos de aportarle creatividad para conseguir despertar el interés del encuestado

Ejemplo: una empresa utilizó una técnica combinada de correo y teléfono en un estudio sobre altos ejecutivos. Se envió a cada uno de los encuestados una pequeña caja de seguridad que contenía el material para la encuesta. Junto a ella se envió una carta en la que se informaba al destinatario de que se le iba a llamar por teléfono, se le daría la combinación de la caja y se le entrevistaría acerca del contenido de la misma.

La entrevista por correo se utiliza en estudios en los que los encuestados tienen bastante interés en el objeto del estudio, como es el caso de los estudios tipo panel.

En estos estudios hay generalmente una muestra pequeña y dispersa.

Se han desarrollado procedimientos para aprovechar lo económico del sistema y, al mismo tiempo, evitar la baja tasa de respuesta. Estos sistemas utilizan encuestados que de antemano han aceptado participar en los estudios. Son los paneles por correo.

4.6.1 MIXTAS COMBINACIONES TELEFÓNICA, CORREO Y PERSONAL

En ocasiones, se pueden reducir costos utilizando el teléfono y / o el correo en partes del estudio.

Por ejemplo, consideremos el caso de una prueba de producto.

1º Mediante entrevista personal seleccionamos a los participantes del estudio y les entregamos la muestra de producto objeto de la prueba.

2º En vez de hacer una segunda visita personal para obtener la evaluación, la podemos obtener por teléfono o por correo, e incluso por la combinación de ambos.

La técnica de la entrevista personal se utiliza generalmente cuando se tiene que mostrar o dar algo a los encuestados, o en estudios complejos, sobre todo de opinión y actitud, así como cuando la duración de la entrevista o el tipo de preguntas formuladas hagan inviables los otros medios.

4.7 ENCUESTAS ESTRUCTURADAS POR SUSCRIPCIÓN

En la actualidad, en el mercado español existen los siguientes tipos de encuestas estructuradas por suscripción:

1. ENCUESTA SECTORIAL DE BIENES DE CONSUMO DURADERO
2. ENCUESTAS ÓMNIBUS
3. ENCUESTAS PANEL

4.8 ENCUESTA SECTORIAL DE BIENES DE CONSUMO DURADERO

En España, y con carácter periódico, diferentes sectores productivos dedicados a la fabricación de bienes de consumo duradero, realizan encuestas para obtener información del mercado. Realizan estos estudios sectores como los del automóvil, electrodomésticos, menaje de cocina, muebles, etc. Cada sector realiza su encuesta. La información que se obtiene es de tipo general para todo el sector, es decir, no incluye aspectos específicos de políticas de Marketing de las diferentes marcas.

La metodología, por lo general, es :

Muestra variable

Tipo de muestreo: estratificado (región/hábitat)

Afijación de la muestra

Periodicidad (anual o semestral) en función del carácter reflexivo de la compra

La información que facilita este tipo de encuestas es del tipo siguiente:

- Parque nacional de artefactos
- Grado de saturación, tanto por ciento de hogares o individuos poseedores
- Grado de penetración, tanto por ciento de artefactos en uso
- Compras anuales
- Compras por modelos
- Estacionalidad de la compraventa
- Compra por primera vez

- Compra por reemplazo
- Forma de compra o adquisición (propia o regalo)
- Lugar de compra. Tipo de establecimiento
- Hábitos de uso
- Actitud hacia el producto
- Grado de conocimiento
- Grado de aceptación
- Grado de preferencia
- Características del entrevistado: sexo, edad, profesión, clase social, nivel de ingresos, número de personas en el hogar, hábitat, región geográfica o Nielsen

Este tipo de encuesta permite a la Dirección de Marketing obtener datos como:

- Participación en compras
- Participación en el mercado
- Participación por segmentos de mercado
- Grado de fidelidad
- Posicionamiento en los diferentes segmentos

4.9 ENCUESTA ÓMNIBUS

Es una técnica de recogida de información, mayoritariamente, mediante entrevistas personales o telefónicas, cuya diferencia con las otras encuestas estriba en que su cuestionario está desarrollado para diferentes temas y productos.

El nombre de ómnibus proviene del sentido de compartir, en este caso, la encuesta. Su empleo está muy extendido en todos los países, debido a su reducido precio. Cualquier empresa puede entrar en el ómnibus simplemente contratando algunas preguntas en una encuesta que se lleva a cabo de un modo continuado y a nivel nacional (por lo general, hay regionales y autonómicos). Los ómnibus emplean una muestra representativa, del 5 al 2% de error.

Normalmente, se puede reservar espacio en el cuestionario ómnibus con muy poca antelación (15 días antes del comienzo, e incluso con una semana caso de TEC²).

² TEC Investigación y Marketing Operativo www.tecmarketing.net
Investigación Comercial

Al distribuirse los costes entre varios clientes los precios se reducen considerablemente, siendo el coste mucho menor que para la realización de un estudio *ad hoc*.

Presenta algunas limitaciones, como que las preguntas son preferiblemente cerradas, fáciles de tabular y el tiempo por tema está a su vez limitado (unos 10 minutos). Así mismo, se deben evitar cuestiones que produzcan algún sesgo que pudiese afectar a otras preguntas.

Por lo general, no se aceptan preguntas que evoquen marcas específicas. Las preguntas suelen ser relativas a:

- Conocimiento y recuerdo de marcas
- Experiencia de producto
- Precios
- Tipo de establecimiento donde efectúa la compra
- Frecuencia de compra

El cuestionario ómnibus se aplica mediante entrevista personal o telefónica, empleando muestras aleatorias o de cuotas mediante submuestreos suplementarios. Se diseña una muestra diferente para cada estudio. Está sujeto a dos tipos de error, el de muestreo y el del sesgo producido por el entrevistador.

Es un procedimiento útil y barato para PYMES que no pueden invertir grandes sumas en Investigación de Marketing, ya que el coste del estudio es compartido entre varias empresas y se recogen datos derivados de una muestra grande.

Las principales *limitaciones* que tiene este método de encuesta son:

- No permite utilizar preguntas complejas

No pueden unirse temas que se refieran a muestras distintas. Por ejemplo, si queremos conocer la opinión de los usuarios de diferentes productos, tales como:

crema de afeitar, muestra hombres / tampax, muestra mujeres

- Los bloques de preguntas no pueden ser largos.

Hay que vigilar el orden del cuestionario de manera que unas respuestas no condicionen las de las preguntas siguientes.

4.10 EL PANEL

Panel es un término inglés que se utiliza par designar a un grupo de personas, establecimientos u organizaciones representativas del universo del que fueron

seleccionados, que facilitan información periódica al investigador que ha constituido el panel, sobre diversos aspectos de interés preestablecido.

El panel lo podemos definir como:

Muestra permanente de una población, formada por un grupo de elementos (consumidores, establecimientos, organizaciones) que aceptan prestarse a encuestas periódicas siguiendo un calendario fijado de antemano.

La técnica de panel es una modalidad de la encuesta por suscripción, en la que se realiza una encuesta repetitiva a una muestra fija.

Sus principales ventajas son:

- Cantidad de información. Por lo general, los panelistas son retribuidos, lo que los convierte en buenos colaboradores facilitando información.
- Analíticas. Permite estudiar los cambios de situaciones.
- Coste relativamente bajo. Por regla general el coste de implantar un panel es elevado, pero consiguiendo bastantes suscriptores el precio individual puede resultar interesante.

Los principales inconvenientes están relacionados con la rotación de sus miembros, así como el sesgo que pueden producir los veteranos y los novatos.

La rotación se produce por cansancio y falta de interés; lo que empezó siendo interesante y novedoso para el panelista, se convierte en una obligación, pasando a ser una carga.

Los miembros antiguos, debido a la rutina y al aburrimiento, tienden a considerarse expertos e influyen en los resultados, mientras que los novatos pueden exagerar la información.

4.10.1 PANEL DE CONSUMIDORES

El panel de consumidores es una técnica de recogida de información cuantitativa a través de una muestra permanente de individuos que relaciona sus datos de consumo en un diario de compras.

Un panel de carácter permanente requiere una buena organización, por lo que es difícil de realizar por cualquier empresa. Por ello, algunos institutos de investigación se encargan de realizar este servicio por su cuenta, comercializándolo después a las empresas.

Análisis de la Investigación Cuantitativa

Para facilitar la formación del panel, los institutos de investigación suelen ofrecer estímulos en forma de obsequios, sorteos, retribución económica, etc. El instituto puede eliminar a algunos panelistas, normalmente por retrasos en el envío de los datos; asimismo, también se da el abandono voluntario de algunos miembros del panel. Esto hace que la renovación de la muestra sea continuada.

La precisión del panel se garantiza normalmente empleando muestras de gran tamaño, lo cual nos da como consecuencia errores de muestreo pequeños. También se forma a los panelistas acerca de cómo rellenar los diarios, y éstos no se procesan por parte del instituto de investigación, hasta transcurrido un período de prueba, cuando se tiene la garantía de su adecuada cumplimentación por parte del panelista.

La principal ventaja de este tipo de panel es que a través del diario se evitan los errores que se pueden dar en las encuestas por fallos de la memoria, ya que en el diario se van recogiendo todos los datos. Este sistema presenta el inconveniente de los registros de compras de productos consumidos fuera del hogar (helados, cenas, copas, etc.), que con frecuencia se dejan en la memoria y para el día siguiente.

La información recogida por el panel de consumidores varía en función del tipo de panel. En la actualidad existen diferentes tipos de paneles. Los más usuales son:

Paneles más habituales

Nombre	Muestra
Panel de hogares	Amas de casa
Baby panel	Madres con niños menores de dos años
Panel de jóvenes	Jóvenes de ambos sexos, hasta 20 años
Panel de televisión	Individuos con T.V.
Panel de automovilistas	Individuos con coche
Panel de radio	Individuos con radio

4.10.2 PANEL DE HOGARES

Está formado por una muestra de amas de casa. Para su realización se selecciona una muestra representativa de hogares (normalmente para todo el país) y se solicita que rellenen un diario detallando todas las compras e indicando cantidades (en unidades) y peso de cada producto comprado, así como marcas, variedades, precio unitario, envase, lugar de compra, etc.

Análisis de la Investigación Cuantitativa

Una vez recogidos los diarios de compras, el centro investigador recopila toda la información teniendo en cuenta aspectos como: hábitat, unidad familiar, tipo de establecimiento donde se realizó la compra, etc.

Normalmente, la información se registra semanalmente, enviándose a los clientes los resultados mensualmente, aportando datos sobre:

- Tamaño de mercado
- Tendencias del mercado
- Participación de marcas
- Precios medios de compra
- Ofertas especiales
- Tamaño del envase

Este panel de hogares tiene la ventaja de que la información suministrada se puede segmentar de acuerdo a los perfiles de los consumidores que compran una determinada marca y los que no la compran. Estos perfiles incluyen características demográficas y psicográficas, definen el estilo de vida y sirven para fijar el público objetivo, pudiendo de esta manera orientar mejor los esfuerzos de comunicación de las empresas en su ajuste producto / segmento de mercado.

Como el panel realiza, además, un seguimiento individualizado de cada hogar, se pueden detectar aspectos tales como:

- Cambio de los hábitos de compra
- Diagnóstico de fidelidad de marca
- Compradores de repetición
- Cambiadores de marca
- Consumidores frecuentes

Por ejemplo, podemos obtener la participación de una marca en el mercado mediante la siguiente fórmula:

$$P_m = P \times Tr \times F_c$$

Donde P_m representa la Participación de marca P representa la Penetración de Mercado Tr representa la Tasa de repetición de compra F_c representa la Frecuencia de compra

La principal crítica que tiene este panel es que provoca un efecto de condicionamiento. Así, puede ocurrir que consumidores veteranos (más de un año), tienden a mostrar más interés por los precios de los productos, que los que no son miembros, tendiendo a comprar marcas más baratas. Aunque este tipo de sesgo no está bien definido, lo que se suele hacer es ir rotando a su miembros, por ejemplo un 25 % anual de rotación, entre bajas espontáneas y sustituciones.

4.10.3 DUSTBIN-CHEK

Lo podemos denominar como comprobación del recipiente de los desperdicios.

Se trata de una modalidad de panel de hogares. Consiste en depositar los envases y las etiquetas de los productos comprados en unos recipientes especiales que posteriormente son recogidos por personal del instituto de investigación para realizar el correspondiente estudio.

A través de este método se obtiene información de los productos consumidos. Evita la escritura en el diario de todas las compras realizadas. Su limitación es que sólo se puede utilizar para productos convenientemente envasados y etiquetados y que se consuman en el hogar.

4.10.4 PANEL DE AUDIÓMETROS (T.V.)

Este panel recoge información sobre en qué cadena y momento del día está el televisor conectado, por medio de un aparato que instala el instituto de investigación en los hogares que forman la muestra.

Antiguamente, para medir la audiencia de televisión se empleaban encuestas y los televidentes debían ir anotando continuamente los programas y cadenas que veían.

En la actualidad, se emplean aparatos denominados audiómetros. Son aparatos que registran si el T.V. está encendido, así como el momento del día y la hora, sin ocasionar molestia alguna a los componentes del hogar.

Los audiómetros son de diferentes características. Los más normales recogen la información en un cassette que se analiza posteriormente en el instituto; otros están conectados a través del teléfono al instituto de investigación, transmitiendo la información directamente a un ordenador central.

Esta técnica está en continua evolución. Por ejemplo, se ha ensayado con aparatos que determinan qué número de personas se encuentran congregadas ante el televisor.

4.10.5 PANEL DE DETALLISTAS

Este panel también recibe los nombres de auditoría de minoristas y panel de establecimientos.

Esta técnica empezó a funcionar en EE.UU. en el año 1929, merced a Arthur C. Nielsen, extendiéndose posteriormente al resto de los países.

En teoría, se puede organizar un panel de estas características para cualquier tipo de establecimientos. Así, podemos hablar de:

panel de electrodomésticos, panel de joyería, panel de cash and carry (mayoristas), panel droguería perfumería, panel detallistas perfumería, panel tiendas de deporte, etc.

Veamos cómo funciona este panel:

Los datos son recogidos de una muestra fija de tiendas, seleccionada para que su volumen de ventas sea representativo de un universo (país, autonomía, región). La muestra de establecimientos comerciales recoge hipermercados, cadenas de tiendas, grandes almacenes, autoservicios, tiendas tradicionales, etc. En este procedimiento no son los dueños o empleados de la tienda quienes realizan la inspección del almacén, sino personal especializado enviado por el instituto de investigación. De esta forma, se evita todo sesgo por parte de la propiedad del establecimiento.

A estos empleados del instituto se les llama auditores. Estos auditores se presentan en las fechas establecidas y realizan el inventario, normalmente cada dos meses, registrando las ventas, stocks, volúmenes de exposición y precios de cada marca.

Las ventas al consumidor se obtienen de la aplicación de la siguiente ecuación

$$\text{VENTAS} = \text{Stock inicial} + \text{pedidos} - \text{Stock final}$$

Stock inicial = son las existencias de un determinado producto que había en el almacén del establecimiento cuando se hizo la anterior auditoría o cuando comenzó el estudio.

Pedidos = compras realizadas por el establecimiento a su proveedor en el período objeto de estudio.

Stock final = existencias que hay en el almacén de ese producto en el momento de efectuar la auditoría. Serán las existencias iniciales de la siguiente auditoría.

Los datos emanados del panel de detallistas permiten evaluar aspectos tales como:

- Tamaño del mercado y tendencias
- Tanto por ciento de la marca y tendencia

Análisis de la Investigación Cuantitativa

- Participación de la marca (por tipo de establecimiento y zona geográfica)
- Ofertas especiales
- Precios
- Estrategias de comunicación (promociones, campañas de publicidad)

A pesar de que permite a las direcciones de las empresas, suscritas a estos datos, gestionar aspectos tales como la distribución, los stocks, los grados de aceptación comercial, el punto débil de estos estudios es que no ofrecen ninguna información sobre motivos de consumo, personas que consumen los productos y que los compran, o las actitudes que tienen hacia los productos.

Los establecimientos que configuran el panel se comprometen a permitir realizar el trabajo de los auditores a cambio de determinadas compensaciones.

El mayor exponente de este sistema es el Nielsen, que publica sus índices de minoristas a escala mundial.

En España, las empresas que adquieren la información Nielsen (índices Nielsen) reciben periódicamente una información de tipo estándar.

Análisis de la Investigación Cuantitativa

4.11 DIFERENCIAS ENTRE LAS DISTINTAS MODALIDADES DE ENCUESTA ESTRUCTURADA

En el esquema siguiente recogemos las principales diferencias existentes entre las encuestas estructuradas ad hoc, ómnibus y panel.

TIPO	AD HOC	ÓMNIBUS	PANEL
INICIATIVA	Cliente	Instituto de investigación	Instituto de investigación
UNIVERSO	En función de las necesidades del cliente	Fijo. Predeterminado por el Instituto	Fijo. Predeterminado por el Instituto
MUESTRA	Variable	Variable	Fija
ESTUDIO	Transversal	Transversal	Longitudinal
CUESTIONARIO	Diseñado en función de las necesidades del cliente	Fijo en las preguntas de calificación y clasificación. En las cuestiones categóricas diseñado en función de las necesidades del cliente	Fijo a lo largo de períodos temporales de contratación
MODALIDAD	Contratación por estudio	Suscripción por pase	Suscripción temporal

5. INTRODUCCIÓN AL ANÁLISIS DE LOS DATOS

5.1 INTRODUCCIÓN

Una vez recogidos los datos y realizado el trabajo de campo, procederemos a su análisis. El análisis cuantitativo de estudios primarios es factible de tratamiento estadístico, ya que normalmente se recoge la información mediante cuestionarios estructurados, lo cual nos permite aplicar programas estadísticos informatizados para transformar los datos en información.

La información así obtenida es la que nos servirá para obtener conclusiones, que aplicaremos en la correspondiente toma de decisiones.

5.2 FASES DEL PROCESO DE ANÁLISIS DE LOS DATOS

La información cuantitativa se ha recogido mediante cuestionarios estructurados, realizando el trabajo de campo diversos encuestadores. Una vez recogida la información, se procede a su análisis, proceso éste que consta de las siguientes fases:

1. Revisión del trabajo de campo y de los cuestionarios
2. Codificación y tabulación
3. Análisis de cada cuestión o ítem
4. Análisis de los ítems por subgrupos. Cruces de preguntas
5. Estudio de las relaciones entre pares de preguntas
6. Estudio de las relaciones entre todas las preguntas
7. Resultados. Conclusiones. Informe

A continuación vamos a reseñar brevemente cada una de estas fases:

5.2.1 REVISIÓN DEL TRABAJO DE CAMPO Y DE LOS CUESTIONARIOS

En esta fase se busca identificar y corregir posibles errores, así como determinar las posibles fuentes de error en los trabajos anteriores. Por ejemplo, respuestas no legibles, cuestionarios mal rellenados, etc.

En ocasiones se subsanan los errores volviendo a realizar la entrevista. Normalmente en la praxis se desecha el cuestionario.

5.2.2 CODIFICACIÓN Y TABULACIÓN

Una vez depurados los cuestionarios, se procede a la codificación de las diferentes preguntas. La codificación consiste en asignar números a cada una de las posibles respuestas.

Ejemplos:

1.- Supongamos una variable dicotómica tal como el sexo

Sexo:	Hombre	☐	(1)
	Mujer	☐	(0)

2.- Consideremos una múltiple respuesta, tal como la edad medida en intervalos

Edad	18 a 25 años	☐	(1)
	26 a 45 años	☐	(2)
	46 a 65 años	☐	(3)
	más de 65 años	☐	(4)

En el primer ejemplo se trata de una variable nominal que toma el valor 1 cuando es hombre y 0 cuando es mujer. En el segundo, la variable está clasificada en categorías, asignándose un valor numérico a cada una de ellas.

También podemos encontrarnos variables cualitativas como la siguiente:

3.- Para ir al trabajo, ¿qué línea de autobús utiliza?

Línea ☐

Normalmente el código será el valor que nos dé el encuestado. Así, si responde que la línea que utiliza es la 40, el código será 40 (en este ejemplo la variable es cualitativa).

En el caso de variables métricas, el propio valor de la medición equivale al código

4. ¿Cuánto gasta, en pesetas, diariamente en transporte?

Gasto ☐

Normalmente, el código será el valor que nos dé el encuestado. Así, si responde que gasta 500 ptas., el código será 500.

En ocasiones una misma pregunta puede contener dos variables, como la siguiente

Análisis de la Investigación Cuantitativa

5.- Indique la importancia que tienen para usted las siguientes características de un frigorífico y si influyó en su decisión de compra:

	Grado de importancia					Influyó en la compra	
	Nada	Poco	Algo	Bastante	Mucho	Sí	No
	(1)	(2)	(3)	(4)	(5)	(1)	(2)
Capacidad							
Dos cuerpos							
Diseño							
Altura							

De cada característica obtenemos dos variables: grado de importancia e influencia en la compra

Las preguntas de respuesta múltiple se descomponen en distintas variables, por ejemplo:

6.- Antes de comprar un frigorífico se asesora a través de:

1. Familia
2. Amigos y conocidos
3. Vendedor de tienda especialista
4. Folletos publicitarios
5. Otros

El encuestado puede elegir más de una respuesta. En este ejemplo tenemos cinco variables y a cada una de ellas le corresponden dos modalidades: Sí y No.

Por tanto la podemos codificar:

Familia	Sí (1)	No (0)
Amigos y conocidos	Sí (1)	No (0)
Vendedor de tienda especialista	Sí (1)	No (0)
Folletos publicitarios	Sí (1)	No (0)
Otros	Sí (1)	No (0)

Hay paquetes estadísticos que permiten el tratamiento de múltiple respuesta, como el SPSS. En estas circunstancias, lo codificaríamos de forma simple.

Familia	1	Folletos publicitarios	4
Amigos y conocidos	2	Otros	5
Vendedor de tienda especialista	3		

Análisis de la Investigación Cuantitativa

Cuando en el cuestionario se utilizan preguntas abiertas se complica el tratamiento, ya que se obtiene, o puede obtenerse, un gran número de respuestas diferentes. Entonces la pregunta abierta se convierte en cerrada, examinando todas las respuestas obtenidas y clasificándolas en categorías, asignando una a cada respuesta.

Una vez codificados los diversos ítems del cuestionario, éstos se recopilan en una matriz de datos; en ésta, las filas serán las diferentes unidades de muestra y en las columnas tendremos las variables. Se trata de realizar una hoja de cálculo. Su formato será parecido al siguiente:

Supongamos un cuestionario de 10 preguntas

nº	V01	v02	v03	v04	v05	v06	v07	v08	v09	v10
1										
2										
...										
n										

5.2.3 ANÁLISIS DE CADA CUESTIÓN O ÍTEM

Una vez que tenemos todos los datos en un fichero u hoja de cálculo, se comienza el análisis, teniendo en cuenta los siguientes pasos:

Se estudia cada pregunta.

En las variables cualitativas, nominales o atributos, se estudia la distribución de frecuencias, los porcentajes y, como medida de tendencia central, la moda.

En las variables métricas, aparte del estudio de las frecuencias y porcentajes, tomaremos como medidas de tendencia central o posición la moda, la mediana y la media; y como medidas de dispersión, la desviación típica y el coeficiente de variación.

5.2.4 ANÁLISIS DE LOS ÍTEMS POR SUBGRUPOS.

Al realizar el análisis individualizado de las cuestiones, podemos encontrar grupos de población que sean de interés (p. ej. votantes y no votantes, en un estudio electoral). En este caso podemos establecer comparaciones entre los diferentes grupos.

Análisis de la Investigación Cuantitativa

5.2.5 ESTUDIO DE LAS RELACIONES ENTRE PARES DE PREGUNTAS

Se estudian las posibles relaciones entre dos ítems en función del tipo de variable. Resumimos los principales sistemas de medida en el siguiente esquema:

TIPO DE VARIABLE	CUANTITATIVA	CUALITATIVA
CUANTITATIVA	Coeficiente de correlación (Pearson) Regresión Coeficiente alfa (α) de Cronbach	Análisis de varianza (Test F) Test “t” de medias
CUALITATIVA		Tabla de contingencia (Chi cuadrado) Correlación de rangos (Spearman) Coeficiente de Cramer

5.2.6 ESTUDIO DE LAS RELACIONES ENTRE TODAS LAS PREGUNTAS:

Se realiza a través de métodos multivariantes. Existen métodos que explican una o más variables en función de las otras, y métodos descriptivos que estudian las relaciones entre todas.

5.2.7 RESULTADOS. CONCLUSIONES. INFORME

Una vez concluida la investigación, se entrega al cliente el correspondiente informe de resultados en la forma acordada (por escrito, en soporte informático, etc.). Este informe contendrá como mínimo los siguientes apartados:

1. Problema estudiado, objetivos de la investigación e hipótesis de trabajo.
2. Metodología seguida: ficha técnica (universo, muestra y su selección, técnica aplicada, fechas del trabajo de campo, tratamiento estadístico, nivel de confianza, probabilidad y margen de error).
3. Exposición de los resultados, conclusiones y recomendaciones (este último apartado no siempre).

5.3 TABLAS DE DATOS

Para analizar la información recogida a través de encuestas, una vez validadas y codificadas, los datos obtenidos se recogen en tablas. Las tablas normalmente son de formato rectangular y comprenden tantas filas como individuos u observaciones se han realizado; la tabla se corresponderá con el tamaño de la muestra **n** (muestra validada), y tendrá tantas columnas como preguntas, ítems formulados o características medidas.

variables

	v01	v02	v03	v04	v05	v06		...	j	...	vp
1											
2											
...											
....											
I									x_{ij}		
...											
“n”											

El valor de intersección de la fila y la columna es el resultado de la valoración que concede el elemento “**i**” a la característica “**j**” y se representa por x_{ij} .

Si el tamaño de la muestra es “**n**” y el número de preguntas o variables es “**p**”, matemáticamente una tabla de datos es una matriz de “**n**” filas y “**p**” columnas. A cada elemento se le asocia un vector correspondiente a su fila, que recoge la valoración que concede ese elemento a cada una de las características: $V_i = (x_{i1}, x_{i2}, x_{i3}, \dots, x_{ip})$

A cada variable se le asocia un vector correspondiente a su columna, que recoge las valoraciones que conceden todos los elementos a esa característica.

$$V_j = (x_{1j}, x_{2j}, x_{3j}, \dots, x_{np})$$

5.4 TIPOS DE TABLAS

5.4.1 TABLAS CUANTITATIVAS

Recogen el valor que toma para el conjunto de elementos (filas) un grupo de variables cuantitativas. Se trata de números que corresponden al valor real.

Ejemplo:

	Edad	Ingresos	Gasto luz	Vivienda	Tfno.	...
1	40	200.000	4.000	55.000	8.000	...
....						
n						

5.4.2 TABLAS DE DATOS ORDINALES Y PREFERENCIAS

Se utilizan para resumir los resultados de las encuestas en las que se pide a los encuestados que ordenen un conjunto de objetos (marcas) de acuerdo con algún criterio (preferencia)

Ejemplo: marcas

Encuesta	A(1)	B(2)	C(3)	D(4)	E(5)
1	1	5	4	3	2
2	2	3	5	1	4
3	5	4	3	1	2
...	...	...	...		
n	1	3	2	5	4

El elemento x_{25} indica que el encuestado 2 valora la marca E con 4.

5.4.3 TABLAS BINARIAS

Se utilizan para variables dicotómicas. Las respuestas de los encuestados se señalan con 0 y 1 (Sí = 1, No = 0).

Ejemplo: preguntas P1, P2, P3, ... etc.

Encuesta	P1	P2	P3	...
1	1	0	0	...
2	0	1	1	...
3	1	1	1	...
...	...	...	...	...
n	...	...	...	...

El elemento $x_{12} = 0$ nos indica que el encuestado 1 ha respondido “no” a la variable 2.

5.4.4 TABLAS DE MODALIDADES

Las preguntas tienen varias alternativas de respuesta, asignando a cada una un código.

Ejemplo:

P 1. El trabajo que desarrolla:

Me gusta mucho	1
Indiferente	2
No me gusta nada	3

La tabla será

Encuesta	P 1	P ...
1	1	...
2	2	...
...	...	...
n	...	...

El elemento x_{21} significa que el encuestado 2 considera indiferente el trabajo que realiza.

5.4.5 TABLAS DISYUNTIVAS COMPLETAS

Se trata de una variante de la anterior; la pertenencia se representa con 1 y la no pertenencia, con 0.

En el ejemplo anterior, la tabla será:

P.1

Encuesta	1	2	3
1	1	0	0
2	0	1	0
...	...	...	...
n	...	...	...

La suma de los elementos de una fila corresponde al número de preguntas.

5.4.6 TABLAS DE PROXIMIDADES Y DISTANCIAS

Se utilizan cuando estudiamos la semejanza o disparidad entre parejas de objetos. Se trata de tablas simétricas, esto es, $x_{ij} = x_{ji}$

Ejemplo:

Supongamos que comparamos las marcas A, B, C, siendo 1 la mínima semejanza y 10 la máxima.

La tabla correspondiente es:

Marca	A	B	C
A	10	2	7
B	2	10	6
C	7	6	10

5.4.7 TABLAS DE SERIES TEMPORALES

Recogen el valor de un conjunto de variables en diferentes momentos de tiempo. Por ejemplo, si estudiamos la evolución del producto interior bruto a lo largo de los años, obtendremos una tabla del tipo siguiente:

Año	P I B
1.990	100
1.995	107
2.000	115

5.4.8 TABLAS MIXTAS O HETEROGÉNEAS

Están formadas por un conjunto de variables de naturaleza diferente. Incluyen variables métricas y atributos o cualitativas. Son las más habituales en la investigación comercial.

Variables

Encuesta	Edad	Nivel estudios	...
1	45	1	...
2	52	2	...
3	31	4	...
....	...	...	...
n	...	...	...

6. ANÁLISIS DE LOS DATOS

Una vez efectuadas correctamente todas las fases expuestas hasta el momento del proceso de Investigación Comercial pasamos a la siguiente: el análisis de los datos.

El análisis de los datos es la interpretación de los resultados obtenidos en las anteriores etapas de la investigación. Hemos de recordar que por muy excelente que sea este análisis, no se eliminan los posibles errores cometidos en las fases anteriores (sesgos, no respuestas y demás errores no muestrales) e, incluso, puede que éstos se acentúen más en los resultados que obtengamos en esta fase.

La cantidad de procedimientos estadísticos es muy amplia. El investigador deberá conocer el número suficiente de técnicas o herramientas que le permitan sacar utilidad de los datos así como conocer cuando es interesante su aplicación, tanto por el tipo de datos de los que dispone como por el tipo de resultados que le proporcionan. Podemos afirmar que cuantas más herramientas conozca el investigador, más posibilidades tendrá de realizar un buen análisis.

Antes de utilizar una determinada técnica hemos de tener en cuenta los siguientes aspectos:

El objetivo que perseguimos. El análisis de los datos persigue expresar de una forma sencilla y reducida las relaciones entre los datos con referencia al problema objeto de la Investigación Comercial. Estamos buscando respuestas al problema de la Investigación Comercial, no tratando de realizar una tesis doctoral o una investigación exhaustiva.

La escala básica de medida en que hemos obtenido los datos, ya que según se trate de una escala nominal, ordinal, de intervalo o de razón, las posibilidades de tratamiento varían.

El número de variables a analizar simultáneamente. Podemos considerar análisis de una sola variable (univariable), de dos variables (bivariable), o de más de dos variables o análisis multivariable.

En este capítulo vamos a estudiar una serie de técnicas para el análisis de una y dos variables, es decir, del análisis univariable y del bivariable.

6.1 MÉTODOS ESTADÍSTICOS CLÁSICOS

6.1.1 CONCEPTOS BÁSICOS

Vamos a reseñar brevemente algunos conceptos de la estadística aplicables a la Investigación Comercial a través de muestras.

6.1.1.1 TIPOS DE ESTADÍSTICA

Normalmente las técnicas estadísticas se clasifican en tres grupos:

- La estadística **descriptiva** utiliza métodos numéricos y gráficos con el propósito de analizar el comportamiento, resumir y presentar la información contenida en un conjunto de datos, sin pretender generalizar los resultados obtenidos.
- La **estadística inferencial** (inferencia estadística) utiliza los datos de una muestra para hacer estimaciones, predicciones u otras generalizaciones sobre un conjunto de datos más amplio (la población), siendo una ayuda en el proceso de toma de decisiones. Al ser una inferencia (estimación o predicción), está sujeta a cierto error, con lo que se deberá calcular el nivel de confianza, medida de la seguridad con que se efectúa la inferencia.
- La **estadística relacional** utiliza los datos de una muestra para medir el grado de relación existente entre dos o más variables. Al ser también una inferencia estará sujeta a cierto error, con lo que tendrá un determinado nivel de confianza.

6.1.1.2 MEDIR

Medir es asignar números o categorías a objetos, sucesos o casos, siguiendo reglas determinadas. De esta manera, en cada una de las unidades estudiadas se recoge información sobre una serie de características que varían de una unidad a otra. A la serie de características se les denomina variables o atributos. Cada variable toma un determinado valor en cada caso (por ejemplo: edad, estado civil, sexo, etc.).

6.1.1.3 ATRIBUTO

Es el carácter de una población no susceptible de ser medida numéricamente. Las diferentes formas que puede presentar el atributo se denominan *modalidades*, también

llamadas *variables cualitativas* por otros autores. El nombre de variable cualitativa es el más utilizado en la Investigación Comercial.

6.1.1.4 VARIABLE

En un sentido estadísticamente académico, una variable es cualquier carácter de una población o universo susceptible de tomar valores numéricos.

Las variables pueden ser continuas y discretas.

Variable continua es aquella en la que, entre dos valores determinados, siempre puede encontrarse otro valor. Por ejemplo, la edad; entre 46 y 47 años siempre podemos encontrar casos intermedios, por ejemplo 46 años y 5 meses, etc.

Variable discreta es aquella en la que, entre dos valores determinados, no se encuentra ningún otro. Por ejemplo, dos familias de 5 y 6 miembros (es imposible encontrar familias intermedias); número de personas en un hogar, (nunca se encontrarán valores intermedios como 5´4 miembros).

En la Investigación de Marketing muchas variables que son continuas (dinámicas), se discretizan; por ejemplo la edad.

“En la práctica de la investigación comercial se acostumbra a hablar de variables cuantitativas y variables cualitativas o categóricas que corresponden, estas últimas, a lo definido como atributo”.

6.1.2 ESCALAS DE MEDIDA

Para realizar mediciones se establecen escalas. Las escalas básicas más usuales son: nominal, ordinal, de intervalo y de razón.

6.1.2.1 ESCALA NOMINAL

En este tipo de escala se asigna un código o nombre a cada caso. La operación lógica que puede realizarse es el establecimiento de igualdad o desigualdad entre los diferentes casos.

Por ejemplo: la variable sexo, con sus categorías de: masculino y femenino; podemos decir que Pepe=Juan; y Tere $\neq$ José. (= igual. $\neq$ no igual, diferente).

Análisis de la Investigación Cuantitativa

Al medir en escala nominal podemos obtener frecuencias (nº de repeticiones o de casos), porcentajes o frecuencias relativas, y si los resultados se expresan en tanto por uno se denominan proporciones. El porcentaje también se denomina frecuencia relativa.

Ejemplo:

Sexo	Código	Frecuencia	Porcentaje	Proporción
Masculino	1	70	60'87	0'6087
Femenino	0	45	39'13	0'3913
	Total	115	100	1

6.1.2.2 ESCALA ORDINAL

En esta escala, además de las relaciones de igualdad y desigualdad, se establece un orden lógico de categorías, por ejemplo, clase social: alta, media-alta, media-media, media-baja, y baja.

Entendemos de esta manera que un individuo de la clase alta está mejor situado que otro de la clase media-alta, y así sucesivamente.

En la escala ordinal, además de asignar un nombre, código o categoría a cada caso o individuo (igual que se hace en la escala nominal), se establece un orden lógico entre las categorías.

A pesar de establecerse un orden, no puede decirse que la distancia o diferencia entre diferentes categorías sea la misma. Es decir no existe igualdad entre los intervalos.

Ejemplo: clase social (alta, media alta, media, media baja y baja)

No podemos decir que la diferencia entre la clase social alta y la media alta sea la misma que la existente entre la clase media baja y la baja.

Clase social	Código
Alta	1
Media alta	2
Media	3
Media baja	4
Baja	5

6.1.2.3 ESCALA DE INTERVALO

Se establecen las relaciones de igualdad / desigualdad, así como el orden lógico de categorías, y la igualdad de intervalos. Para ello es preciso fijar una unidad de medida constante y uniforme a lo largo de todos los posibles valores de la variable.

El punto cero es un valor consensuado, pero que no tiene un significado de ausencia total del aspecto medido. Por ejemplo: la diferencia entre 21°C y 28°C es la misma que entre 31°C y 38°C.

Es decir, podemos establecer: igualdad y desigualdad $20^{\circ}\text{C} = 20^{\circ}\text{C}$, y $25^{\circ}\text{C} \neq 18^{\circ}\text{C}$.

Existe orden $20^{\circ} > 19^{\circ} > 18^{\circ}$, etc.

Igualdad de intervalo, es decir el aumento de 1°C, representa lo mismo a 20°C que a 4°C.

Y el origen o cero de la escala, (0°C), es un valor consensuado que no significa la ausencia de valor. Recordar el conocido chiste ¡Qué bien estamos a 0°C ! No hace “ni frío ni calor”.

6.1.2.4 ESCALA DE RAZÓN O PROPORCIÓN

Aparte de todas las propiedades de las escalas de intervalo, se establece un cero absoluto, que representa la ausencia total del aspecto que se está valorando. La existencia de este cero absoluto es un requisito necesario para poder realizar comparaciones mediante cocientes. Las operaciones matemáticas que utiliza el hombre con mayor frecuencia para realizar comparaciones son la resta y la división.

Es decir, en las escalas de razón o proporción existen las relaciones de igualdad, orden, igualdad de intervalos y además, el cero absoluto.

6.1.2.4.1 RESUMEN ESCALAS BÁSICAS DE MEDIDA

En la tabla siguiente se resumen las características y relaciones de las diferentes escalas

Escala	Relaciones que la definen
Nominal	Equivalencia
Ordinal	Equivalencia Mayor que
De intervalo	Equivalencia Mayor que Razón conocida entre cualesquiera intervalos. Cero consensuado
De razón	Equivalencia Mayor que Razón conocida entre cualesquiera intervalos Razón conocida entre cualesquiera de dos valores de la escala Cero como ausencia de la característica que se mide

6.1.3 ELEMENTOS DE LA ESTADÍSTICA INFERENCIAL

La estadística inferencial tiene los siguientes elementos característicos:

Población: un determinado conjunto de elementos (personas, objetos, transacciones, sucesos...) de los que queremos estudiar una característica o propiedad (por ejemplo, ¿cuál es el consumo medio de...?, o ¿cuál es el porcentaje de consumidores de...?).

Muestra: un subconjunto de elementos de la población. Son los que realmente estudiamos, de los que obtenemos información.

Estimador o inferencia estadística: una estimación, predicción u otra generalización sobre la característica estudiada de la población, basada en la información contenida en la muestra.

Nivel de confianza: una medición de la seguridad con que se efectúa la inferencia (estimación o predicción).

6.2 ANÁLISIS UNIVARIABLE

Vamos a reseñar esquemáticamente las técnicas estadísticas utilizadas en el estudio de una sola variable.

6.2.1 DESCRIPCIÓN DE LOS DATOS

Para obtener una descripción de los datos se realizan básicamente las siguientes mediciones:

1. Frecuencias, absolutas y relativas
2. Medidas de tendencia central
3. Medidas de dispersión
4. Medidas en cuanto a la forma de distribución, coeficientes de asimetría y apuntamiento o curtosis
5. Otras mediciones como ratios o relaciones por cociente, normalización de la variable y medidas de concentración tales como curva de Lorenz, índice de Gini o medial o mediana.

A continuación, resumimos las cuatro primeras de estas medidas³:

6.2.2 FRECUENCIA

Es el número de veces que se presenta un determinado valor de una variable. Se distinguen dos tipos de frecuencias.

A la que nos viene indicada por el número de veces o casos que presenta una variable objeto de estudio, se le denomina frecuencia absoluta o simplemente frecuencia.

Si calculamos el porcentaje (porcentaje es la proporción en que contestaron o eligieron una modalidad, por cien), se denomina frecuencia relativa.

En la tabla siguiente resumimos las principales características

³Para las muestras utilizaremos letras latinas, mientras que cuando nos refiramos al universo se utilizarán griegas.

TABLA: CONCEPTOS RELACIONADOS CON LA DISTRIBUCIÓN DE FRECUENCIAS

Frecuencia absoluta (*frequency*): número de veces que se repite cada valor en el conjunto de los datos.

Porcentaje o frecuencia relativa (*percent*): frecuencia absoluta dividida por el número total de observaciones. Se suele expresar en tantos por ciento. Por tanto, es el porcentaje de repeticiones de un determinado valor en el total de la muestra, pero incluye el porcentaje correspondiente a los valores perdidos (*missing cases*).

Porcentaje válido (*valid percent*): frecuencia absoluta dividida por el número de observaciones con información válida. Se excluyen del cálculo los valores perdidos, por lo que indica el porcentaje de respuestas sobre el total de la información conocida.

Porcentaje acumulado (*cum percent*): porcentaje de observaciones que están por encima o por debajo de cierto valor.

Intervalos de clase: grupo de observaciones con valores semejantes. El intervalo queda definido por un valor mínimo, un valor máximo y una marca de clase (el valor central). Por regla general, todas las clases que se definen tienen la misma amplitud (distancia entre el mínimo y el máximo de cada clase).

Frecuencias absolutas y relativas de un intervalo de clase: número o porcentaje de observaciones con un valor que pertenece al intervalo de clase.

Si el valor máximo y el valor mínimo de dos clases consecutivas coinciden, se considera que dicho valor concreto es de la clase superior (donde es mínimo).

6.2.3 MEDIDAS DE TENDENCIA CENTRAL

También se denominan medidas de posición. Son las que resumen la información aportada por todos los elementos objeto de estudio en un valor central. La tendencia central nos da una descripción concisa del promedio o funcionamiento característico del grupo como un todo.

Las medidas de la tendencia central nos permiten comparar dos o más grupos en función de su funcionamiento característico.

Existen varias medidas de tendencia central. Su elección se basa en el tipo de escala que utilicemos. Recordemos que estas escalas son las siguientes: métrica, ordinal y nominal.

Investigación Comercial

Toda la información medida en una escala superior, se puede convertir en otra inferior.

Las medidas de posición o de tendencia central son las siguientes:

La media, la moda y la mediana.

La *media* es el promedio aritmético de los valores obtenidos por la variable.

$$m = \sum_{i=1}^n \frac{x_i}{n}$$

Donde m es la media de la muestra, n es el tamaño de la muestra y x_i cada valor que toma la variable “ i ”

La *moda* es el valor o intervalo de la variable que más se repite, el valor más frecuente.

La *mediana* es el valor de la variable que deja, por encima y por debajo, el 50% de los casos. En la práctica es difícil encontrar un valor de la variable que parta la distribución de frecuencias en dos. Por ello, formalmente, la mediana es el mínimo de los valores tal que *la frecuencia acumulada es igual o mayor al 50%*. Los datos deben estar ordenados en orden ascendente o descendente.

Consideraciones generales acerca de las medidas de posición o de tendencia central:

La moda es la única medida de posición que podemos aplicar a los datos medidos en escala nominal.

Para poder calcular la mediana se necesita, al menos, disponer de datos medidos en una escala ordinal, no siendo posible su cálculo en el caso de datos medidos en escala de tipo nominal.

En el cálculo de la media se precisa disponer, al menos, de una unidad de medida. Esto sólo lo cumplen las métricas, es decir, las escalas de intervalo y las de proporción.

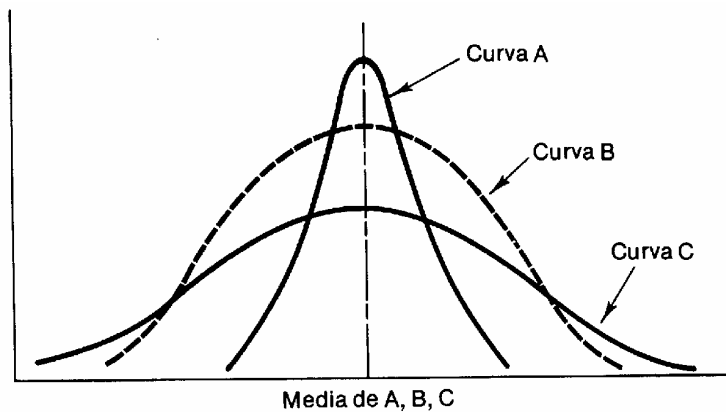
6.2.4 MEDIDAS DE DISPERSIÓN

Las medidas de tendencia central o de posición, aunque son muy útiles para resumir la información contenida en una variable, no aportan toda la realidad de la misma. Con el fin de dar solución a este problema se diseñaron las medidas de dispersión.

Las medidas de dispersión están diseñadas para dar información sobre el error que se comete al considerar todos los casos iguales a la tendencia central.

Existen muchas posibilidades de calcular la dispersión de una variable. Las medidas de dispersión más usuales son: amplitud o rango, la desviación intercuartil, la varianza, la desviación típica y el coeficiente de variación.

Las medidas de tendencia central, la media, la mediana y la moda, sólo facilitan una parte de la información contenida en los datos. Para conocer mejor el comportamiento de las observaciones es conveniente conocer también su **dispersión o variabilidad** (cómo varían los valores de la variable o hasta qué punto son todos muy parecidos o muy diferentes).



6.2.4.1 AMPLITUD O RANGO

Es la diferencia entre el valor más pequeño y el más alto de la variable en estudio. Por ejemplo, si tenemos los siguientes datos para una variable

5, 7, 8, 9, 15, 17, la amplitud sería $17 - 5 = 12$

La amplitud también se llama recorrido; es una buena medida de dispersión por su claro significado, dependiendo sólo de los dos valores extremos. En el caso de una variable muy concentrada, pero con dos extremos muy diferenciados, obtenemos una medida de dispersión poco representativa al calcular la amplitud.

6.2.4.2 RECORRIDO INTERCUARTÍLICO

Se define como la diferencia entre el tercer y el primer cuartil.

Viene dado por la siguiente fórmula:

$$R_I = Q_3 - Q_1$$

6.2.4.3 LA DESVIACIÓN INTERCUARTIL

Es la diferencia entre los dos valores que ocupan los percentiles 25% y 75% dividida por dos.

Tabla: RESUMEN MEDIDAS ELEMENTALES DE DISPERSIÓN

Mínimo (*mínimum*): es el valor más bajo de las observaciones.

Máximo (*máximum*): es el valor más alto de las observaciones.

Rango o amplitud (*range*): es la diferencia ente el máximo y el mínimo.

FRACTILES

Fractil: valor por debajo del cual se encuentra una fracción o proporción de los datos.

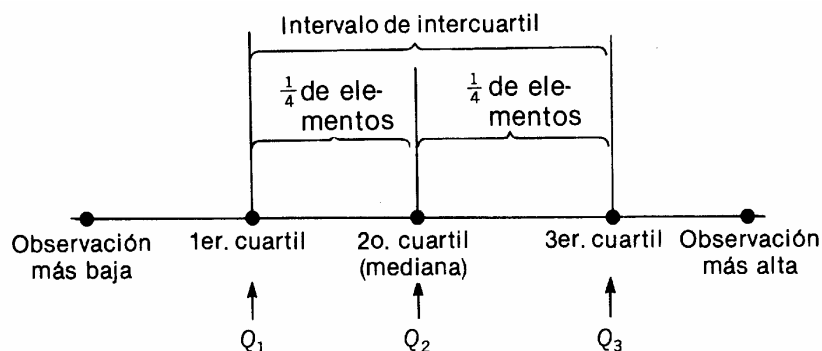
Percentil: la variable se ordena de menor a mayor, entonces se divide en 100 partes iguales (100 grupos con el mismo número de observaciones) y se estudia el valor que alcanza en cada parte. Ejemplo: el percentil al 50% es la mediana.

Cuartil: los datos se dividen en 4 partes iguales, obteniendo:

$$Q_1, Q_2, Q_3 \text{ y } Q_4$$

(Q2 es la mediana y Q4 coincide con el máximo).

Rango o amplitud intercuartil: diferencia entre Q3 y Q1.



Entre el primer y tercer cuartil, en la amplitud intercuartil, se encuentran el 50% de las observaciones centrales.

6.2.4.4 DESVIACIÓN MEDIA

Se define como la media aritmética de las desviaciones entre los valores de la variable y la media aritmética, en valor absoluto. Su fórmula es:

$$DM = \frac{\sum_{i=1}^n |x_i - m|}{n}$$

Si todas las observaciones fueran iguales, la DM sería cero. Cuanto mayor sea el valor de la DM mayor será el grado de dispersión.

Esta medida es poco utilizada debido a que no se puede manipular algebraicamente.

6.2.4.5 VARIANZA Y DESVIACIÓN TÍPICA O ESTÁNDAR

La varianza es una medida de dispersión de variables métricas.

Si pretendemos evaluar exactamente cual es el error que cometemos al asignar el valor de la media a todos los casos de la distribución de la variable, hay que calcular el promedio de las diferencias con la media de todos los valores.

Para poder hacer un promedio, necesitamos una unidad de medida; por consiguiente, la variable deberá estar, al menos, en escala de intervalo.

La varianza es el promedio de las diferentes desviaciones de variable respecto de la media al cuadrado. Se representa mediante S^2 ó σ^2 , según se refiera a la muestra o al universo. Su fórmula es:

$$s^2 = \frac{\sum_{i=1}^n (x_i - m)^2}{n}$$

La desviación típica es la raíz cuadrada de la varianza, el error promedio que cometemos al asignar la media a cada caso. La desviación típica es la distancia de los distintos valores de la variable a la media.

Cuanto menor es la desviación típica, más representativa y más próxima a la realidad es la representación de la variable mediante la media.

La desviación típica presenta la siguiente propiedad:

La desviación típica nos permite realizar rápidamente cálculos sobre entre qué valores se encuentran el 95,5% de los casos, en las variables que siguen una distribución normal (la mayor parte de la variables sigue este tipo de distribución). Para realizar el cálculo,

Análisis de la Investigación Cuantitativa

hay que sumar y restar dos veces la desviación típica al valor de la media. Se representa por:

$$m \pm 2s \quad \text{o bien} \quad \mu \pm 2s$$

Según nos refiramos a la muestra o al universo o población.

TABLA: CONCEPTOS VARIANZA Y DESVIACIÓN TÍPICA

Varianza (variance) de la población: suma de cuadrados de las distancias entre la media y cada elemento, dividido por el número total de observaciones de la población. Es una distancia promedio a la media.

$$\sigma^2 = \frac{\sum (x - \mu)^2}{N}$$

Varianza de la muestra: es equivalente a la de la población, pero se divide por el número de observaciones de la muestra menos 1 (ya que la media de la muestra es un dato conocido):

$$s^2 = \frac{\sum (x - \bar{x})^2}{n - 1}$$

Desviación estándar o típica (Std deviation): es la raíz cuadrada positiva de la varianza, ya sea de la población o de la muestra.

Está medida en las mismas unidades que la variable estudiada.

6.2.4.6 COEFICIENTE DE VARIACIÓN

Es una medida de dispersión relativa, que se utiliza para comparar variables con diferentes unidades de medida.

Consiste en dividir la desviación típica por la media y después expresarla en porcentaje.

Este coeficiente está exento de unidades.

Investigación Comercial

Su fórmula es:

$$CV = \frac{s}{m} \times 100$$

Tiene que ser $m \neq 0$.

El coeficiente de variación tiene la ventaja de que pierde las unidades de medida, por lo que puede servir para comparar las dispersiones relativas de dos variables que tienen diferente unidad de medida.

6.2.5 MEDIDAS RELATIVAS A LA FORMA DE LA DISTRIBUCIÓN

6.2.5.1 DISTRIBUCIÓN

Una distribución es una serie de valores separados, colocados u ordenados según una magnitud. Un conjunto de códigos ordenados y sus frecuencias correspondientes se llaman “distribución de frecuencias”.

6.2.5.2 COEFICIENTE DE SESGO O ASIMETRÍA (SKEWNESS)

Con las medidas de asimetría se intenta medir si las observaciones están dispuestas simétrica o asimétricamente respecto a un valor central, generalmente la media aritmética.

El coeficiente más utilizado es el de R. A. Fisher, que viene expresado por la siguiente fórmula:

$$CA = \frac{\sum_{i=1}^n (x_i - m)^3}{ns^3}$$

Se dan las siguientes alternativas:

$CA = 0$ se trata de una distribución simétrica.

$CA > 0$ es asimétrica a la derecha. Positiva.

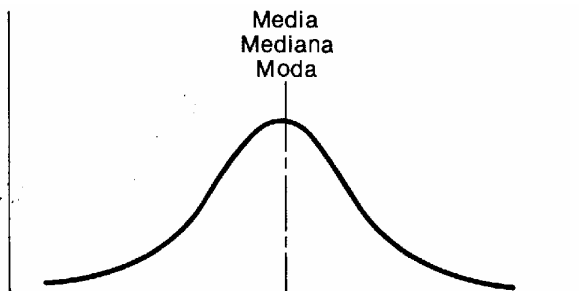
$CA < 0$ es asimétrica a la izquierda. Negativa.

Las distribuciones **simétricas o insesgadas** son aquellas que dejan el mismo número de observaciones a la izquierda que a la derecha de la media.

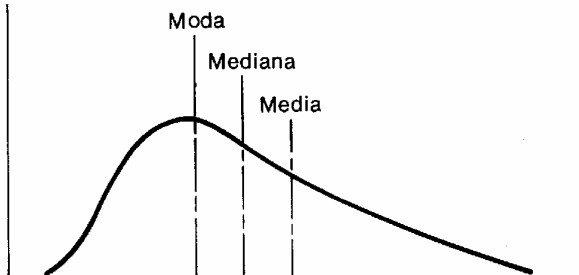
La simetría se mide a través del **sesgo (skewness)**.

En la gráfica siguiente exponemos las correspondientes situaciones⁴

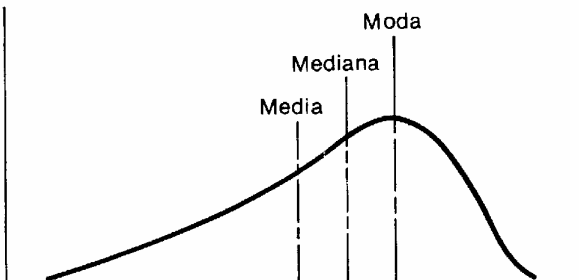
Por tanto, cuando la distribución es simétrica, sesgo = 0, la media es igual a la mediana (e iguales a la moda, si sólo existe una):



Si la mediana es más pequeña que la media la distribución está sesgada a la derecha (hacia valores grandes), sesgo > 0:



Si la mediana es más grande que la media la distribución está sesgada a la izquierda (hacia valores pequeños), sesgo < 0:



6.2.5.3 COEFICIENTE DE CURTOSIS O APUNTAMIENTO

Las medidas de apuntamiento o curtosis son aplicables a distribuciones campaniformes, unimodales simétricas (no en U) o con una ligera asimetría. Se mide la mayor o menor concentración de las frecuencias alrededor de la media, es decir, su nivel de apuntamiento. Se toma como distribución tipo la distribución normal. La distribución normal cumple que:

$$m_4 = 3 s^4 \quad s = 1 \text{ y } m = 0$$

Donde **m_4** es el momento de orden 4 respecto de la media o momento central que viene dado por la expresión siguiente:

⁴ Universidad Autónoma Barcelona. Apuntes de Cátedra Investigación Comercial 1 Teresa Obis
Investigación Comercial

$$m_4 = \frac{\sum_{i=1}^n (x_i - m)^4}{n}$$

El coeficiente de curtosis viene dado por la siguiente expresión:

$$CK = \frac{m_4}{s^4} = \frac{\sum_{i=1}^n (x_i - m)^4}{ns^4}$$

Como CK para la distribución normal es igual a 3 se suele usar el coeficiente corregido, también llamado coeficiente de exceso y que es:

$$CK_2 = \frac{\sum_{i=1}^n (x_i - m)^4}{ns^4} - 3$$

Los valores que toma son:

$CK_2 = 0$ es decir $CK = 3$ Es una distribución mesocúrtica, sin exceso.

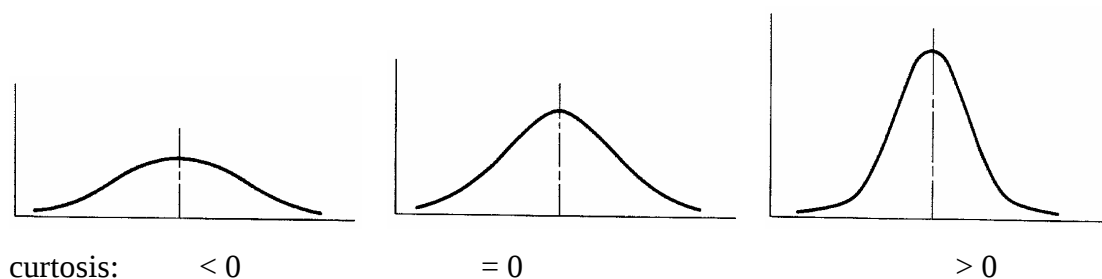
$CK_2 > 0$ es decir $CK > 3$ Distribución leptocúrtica. Con exceso. Puntigrada.

$CK_2 < 0$ es decir $CK < 3$ Distribución platicúrtica. Es decir achatada

En el gráfico siguiente se pueden observar las diferentes formas de la “curtosis” o apuntamiento

Curtosis: grado de apuntamiento (pico) de la distribución. Una distribución en "forma de campana" tiene una curtosis igual a 0.

Si está más concentrada (puntiaguda) la curtosis es > 0 ; si es más plana (con mayor dispersión) la curtosis es < 0 .



6.2.6 ¿CÓMO REALIZAR INFERENCIAS?

Cuando lo que pretendemos es comparar los valores obtenidos en la investigación con otros predeterminados o bien realizar inferencias, lo que se hace es aplicar determinadas pruebas o test en función de la escala de medida de la variable objeto de estudio. Lo podemos resumir en la siguiente tabla:

Escala	Frecuencia	Tendencia central	Dispersión	Prueba
Nominal	Absoluta, relativa	Moda		Chi cuadrado, binomial
Ordinal	Absoluta, relativa, acumulación, percentiles	Moda, Mediana	Desviación intercuartil	
Intervalo	Absoluta, relativa, acumulación, percentiles	Moda, Mediana Media ⁵	Desviación intercuartil Varianza, desviación típica	
Proporción razón	Absoluta, relativa, acumulación, percentiles	Moda, Mediana Media	Desviación intercuartil Varianza, desviación típica Coeficiente de variación	

⁵ En ocasiones no tiene ningún significado, debido al desconocimiento de los intervalos de clase con exactitud.

7. CONTRASTE DE HIPÓTESIS

7.1 CONCEPTOS GENERALES

7.1.1 HIPÓTESIS

La palabra “hipótesis” tiene su origen en los términos griegos *thesis*, que significa “lo que se pone”, e *hipo*, que equivale a “debajo”. Hipótesis es, por tanto, lo que se pone debajo, es decir lo que se supone.

Desde el punto de vista de la investigación comercial, las hipótesis se pueden definir como soluciones probables previamente seleccionadas al problema planteado, que se tienen que confirmar a través del proceso de investigación con los hechos.

La primera prueba a la que debemos someter una idea nueva es a la de su coherencia. Una vez superado este primer paso, la idea o conjetura pasa a denominarse hipótesis y deberemos contrastarla a través de la denominada prueba de hipótesis.

7.1.2 PRUEBAS DE HIPÓTESIS

Para K. R. Popper en su obra *El desarrollo del conocimiento científico* (Ediciones Paidós, 1967), una hipótesis es “una conjetura expresada en términos de alto contenido informativo”.

Para que la hipótesis se convierta en teoría deberá ser validada mediante pruebas de hipótesis. Estas pruebas actúan como un filtro, eliminando las hipótesis falsas y manteniendo las verdaderas. Podemos afirmar que cuantas más pruebas haya superado una hipótesis, más convencimiento tendremos acerca de su veracidad.

En la praxis se ha escogido que el nivel de filtro sea del 95%, es decir, se eliminan el 95% de las hipótesis falsas.

En el campo científico es de gran importancia la repetición de la prueba de hipótesis por diferentes equipos de investigadores y en diferentes partes del mundo. Esto significa que cuantas más oportunidades se han dado a una hipótesis de ser abandonada y más veces ha superado las pruebas, más seguros estaremos de su veracidad.

En el campo de la Investigación Comercial hemos de tener presente que trabajamos con muestras, que los estadísticos son aleatorios y que, para el estadístico, muestras distintas pueden generar diferentes valores.

Una hipótesis estadística es una afirmación con respecto a alguna característica de interés. En las pruebas de hipótesis hay que depurar los resultados mediante estudios más profundos.

Contrastar una hipótesis es decidir si se rechazan o no unos determinados supuestos que planteamos acerca de un universo o población, partiendo de los datos obtenidos con una muestra, midiendo a la vez el riesgo de error correspondiente a cada una de las posibles decisiones.

Para realizar la prueba de hipótesis tenemos que definir previamente la hipótesis nula (H_0) y la hipótesis alternativa (H_1)

7.1.2.1 HIPÓTESIS NULA

Es la hipótesis que se realiza sobre un determinado fenómeno con el propósito de comprobarla o rechazarla. Se representa por H_0 .

7.1.2.2 HIPÓTESIS ALTERNATIVA

Cualquier otra hipótesis que se formule diferente a la hipótesis nula. La hipótesis alternativa se representa por H_1 .

Ejemplo:

Partimos de dos supermercados, A y B; queremos demostrar que la diferencia de sus ventas es debida a la diferencia de permanencia de los clientes en el establecimiento.

Para ello, medimos el tiempo medio de estancia en la tienda de los clientes de A y B.

Definiremos las siguientes hipótesis:

H_0 : las medias de estancia, de los clientes, en la tienda son iguales $m_A = m_B$

H_1 : las medias de estancia en la tienda son diferentes $m_A \neq m_B$

7.1.3 CARACTERÍSTICAS Y METODOLOGÍA

Las principales características de las hipótesis nula y alternativa son:

- Una y sólo una de las dos puede ser cierta

Análisis de la Investigación Cuantitativa

- Mientras H_0 indica una única posibilidad, la hipótesis alternativa comprende infinitas
- Lo que creemos y tratamos de demostrar es la hipótesis alternativa
- Si podemos sospechar que la hipótesis nula no es cierta, entonces obligatoriamente deberá ser cierta la hipótesis alternativa

La metodología para realizar un contraste de hipótesis consiste en la realización de los siguientes pasos:

- Definir H_0 y H_1
- Definir un parámetro de medida que relacione los parámetros muestrales y poblacionales y la función de probabilidad
- Establecer un criterio para juzgar si el valor calculado del parámetro es compatible con H_0
- Realizar el muestreo
- Obtener conclusiones de conformidad con lo expuesto

Con el fin de comprender mejor los diferentes conceptos, vamos a realizar un ejemplo sencillo y clásico.

Para nuestro estudio, utilizaremos una moneda cargada.

Las correspondientes hipótesis serían:

Hipótesis nula (H_0). La probabilidad de que salga cara es 0'5, es decir, la moneda no está cargada $H_0 : p = 0'5$.

Hipótesis alternativa (H_1). La moneda esta cargada. Esto es, $H_1 : p \neq 0'5$.

Para elegir una de las dos hipótesis lanzamos la moneda al aire 100 veces.

Si H_0 fuera cierta el resultado que obtendríamos sería 50 veces cara y 50 veces cruz.

Supongamos que obtenemos cara en 51 ocasiones y cruz en 49. Esta pequeña diferencia sería totalmente explicable por el azar o la suerte. Por consiguiente, nuestra conclusión sería que no hemos encontrado nada que nos haga sospechar que la moneda esté cargada, luego no podríamos decir que la hipótesis nula fuese falsa (no podríamos rechazar H_0).

Supongamos que al lanzar 100 veces la moneda, obtenemos 97 veces cara y 3 cruz. La diferencia es lo suficientemente significativa como para hacernos dudar de la validez de

la hipótesis nula y afirmaríamos que la moneda está cargada o bien que abandonamos la hipótesis nula H_0 (se puede rechazar la H_0).

Pero supongamos que después de 100 lanzamientos de la moneda, hemos obtenido 68 caras y 32 cruces. En esta situación unas personas no rechazarían la H_0 , considerando que la desviación es fruto del azar, mientras que otras sí que rechazarían la H_0 .

Con el fin de dar solución a este tipo de situaciones se establece un criterio homogeneizador. La estadística propone como criterio rechazar la hipótesis nula H_0 cuando la probabilidad de un resultado tanto más extremo, más pequeña que un cierto valor escogido *a priori*. Basándose en la experiencia se acepta que ese valor sea 0'05.

Una vez aceptado ese valor deberemos calcular las probabilidades de los diferentes resultados muestrales suponiendo cierta la H_0 .

Teniendo en cuenta que el producto de $p * n$ es mayor que 5, la proporción p obtenible en las infinitas muestras sigue una distribución normal de media p y desviación típica

$$S_p = \sqrt{\frac{p(1-p)}{n}}$$

siendo n el número de casos de la muestra.

En el ejemplo tenemos que H_0 nos dice que $p = 0'5$. Si lanzamos la moneda 100 veces tendríamos que

$$S_p = \sqrt{\frac{0'5(1-0'5)}{100}} = 0'05$$

Se simboliza como $p \rightarrow n(0'5, 0'05)$

Esto significa que la proporción de caras obtenidas en una muestra de 100 casos sigue una distribución normal de media 0'5 y desviación típica 0'05.

Siguiendo las propiedades de la distribución normal vamos a calcular entre qué dos valores se situarán el 95% de los resultados obtenidos con las infinitas muestras de 100 lanzamientos de la moneda.

Si recordamos que para el valor del 95% la razón crítica $Z = 1'96$. Aplicando la propiedad de la distribución normal $\mu \pm Z\sigma$ obtenemos

$0'5 \pm 1'96 * 0'05$, lo cual nos conduce a los siguientes valores $p_1 = 0'402$ y $p_2 = 0'598$

Por consiguiente, si es cierta la H_0 , la probabilidad de encontrar una proporción muestral de caras entre 0'402 y 0'598 será igual al 95% (0'95).

La probabilidad de encontrar valores fuera de ese intervalo será de 0'05.

De conformidad con lo expuesto, estableceríamos un sistema de decisión del tipo siguiente:

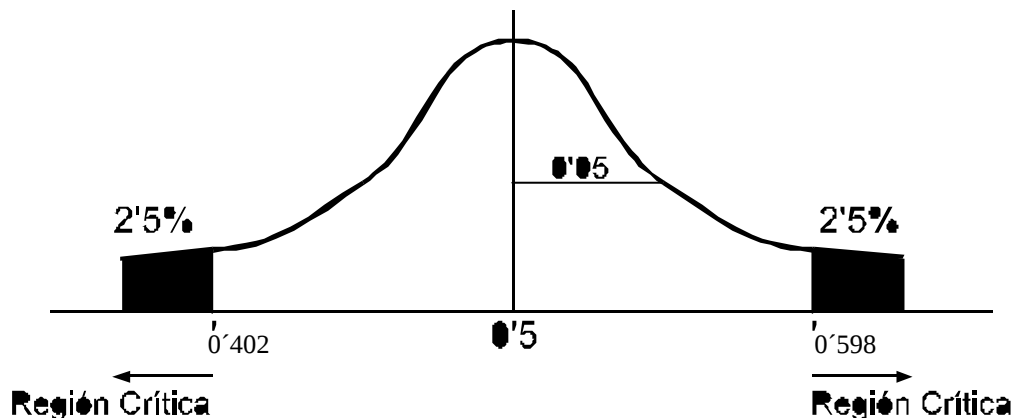
Si la proporción de caras, al lanzar 100 veces la moneda, se sitúa entre 40'2 y 59'8%, la H_0 será cierta.

Si la proporción se sitúa por debajo del 40'2 y por encima del 59'8, se dirá que la probabilidad de observar estos resultados es tan pequeña (0'05) como para sospechar que la H_0 es falsa.

A la zona extrema y alejada, que tiene una pequeña probabilidad de ocurrir si es cierta la H_0 y que si se observa un resultado en ella nos permite rechazar la hipótesis nula, se la denomina *región crítica*.

La representación gráfica para nuestro ejemplo es la siguiente:

Gráfica



Esquema de decisión

Si el resultado se sitúa en la región crítica se rechaza la hipótesis nula H_0

Si el resultado no se sitúa en la región crítica, no se rechaza la hipótesis nula H_0 , rechazando entonces la hipótesis alternativa H_1 .

7.2 TEST DE HIPÓTESIS

También se denomina test de significación o contraste. Se utiliza para designar el procedimiento que se utiliza para contrastar la validez de una hipótesis. Son pruebas estadísticas que se utilizan para determinar si los resultados obtenidos con una muestra o dos elegidas al azar difieren marcadamente de aquellos que habría que esperar con la hipótesis planteada y la variación debida al muestreo.

7.2.1 OBJETIVO

El objetivo de este tipo de test es:

Observar si la variable estudiada se comporta aleatoriamente

Verificar si la media de la muestra estudiada pertenece a la media del universo estudiado

Comprobar si la proporción obtenida en el estudio muestral es la misma que la de la población

Ver si los datos obtenidos en la muestra corresponden a algún tipo de distribución conocida

Observar si los valores obtenidos siguen unos patrones esperados

7.3 METODOLOGÍA DEL TEST DE HIPÓTESIS

Tal y como indicamos en el punto anterior, el test de hipótesis es una forma de decidir, de manera objetiva, si los resultados obtenidos a través de la Investigación Comercial muestran una realidad o son simplemente una consecuencia de la aleatoriedad de la muestra.

En este proceso se diferencian las siguientes etapas:

1. Formulación de las hipótesis
2. Elección del nivel de significación
3. Elección de la prueba o test
4. Interpretación

7.3.1 FORMULACIÓN DE LAS HIPÓTESIS

Desde el punto de vista de la investigación comercial, las hipótesis se pueden definir como soluciones probables previamente seleccionadas al problema planteado, que se tienen que confirmar con los hechos en el proceso de investigación.

Una hipótesis es, por tanto, una proposición relativa a un problema, a la forma o a los parámetros de una distribución. Cuando se refiere a un valor concreto para un parámetro, se denomina hipótesis simple. Si se refiere a un intervalo para un parámetro, hipótesis compuesta.

En el campo de la Investigación Comercial hemos de tener presente que trabajamos con muestras, y que los estadísticos son aleatorios y que muestras distintas pueden generar diferentes valores para el estadístico.

Una hipótesis estadística es una afirmación con respecto a alguna característica de interés. En las pruebas de hipótesis se trata de depurar los resultados mediante estudios más profundos.

Contrastar una hipótesis es decidir si se rechazan o no como ciertos los supuestos que planteamos acerca de un universo o población, partiendo de los datos obtenidos con una muestra, midiendo a la vez el riesgo de error correspondiente a cada una de las posibles decisiones.

El proceso de la prueba o test de hipótesis comienza definiendo la hipótesis, que se confirmará o rechazará de conforme al resultado que obtengamos en el test.

Para realizar la prueba de hipótesis tenemos que definir previamente la hipótesis nula H_0 y la hipótesis alternativa H_1 que, recordemos de nuevo, son:

Hipótesis nula. Hipótesis que se realiza sobre un determinado fenómeno con el propósito de comprobarla o rechazarla⁶. Se representa por H_0 . Se refiere a determinadas características de la población, denota ausencia de diferencias y es la que se mantendrá en el caso de que los resultados de la prueba o test no muestren falsedad.

⁶Siendo estrictos en vez de decir se acepta la hipótesis nula debemos decir “no se rechaza la hipótesis nula”

Hipótesis alternativa. Cualquier otra hipótesis que se formule diferente a la hipótesis nula. Se representa por H_1 . Denota la existencia de diferencia; es una forma de negación de la H_0 .

La formulación de ambas hipótesis debe hacerse de forma muy rigurosa y precisa.

Las principales características de las hipótesis nulas y alternativas son:

1. Una y solamente una de las dos puede ser cierta
2. Mientras H_0 indica una única posibilidad, la hipótesis alternativa comprende infinitas
3. Lo que creemos y tratamos de demostrar es la hipótesis alternativa
4. Si podemos sospechar que la hipótesis nula no es cierta, entonces obligatoriamente deberá ser cierta la hipótesis alternativa

En función de la H_0 se van a definir los valores esperados (teóricos). La comparación entre estos valores y los observados nos proporcionaran la decisión de acuerdo con las normas preestablecidas.

La hipótesis se puede enunciar de forma unidireccional (de una cola), si el sentido o dirección de la hipótesis alternativa es conocido. Ejemplo. la cuota de audiencia de una cadena local de TV es superior al 25%. Las hipótesis las podemos formular

$$H_0 : P \leq 25\% \text{ y la alternativa será } H_1 : P > 25\%$$

También se pueden formular de forma bidireccional, bilateral o test de dos colas (Two-tailed test), si la hipótesis alternativa puede tomar cualquier sentido. Siguiendo con el ejemplo anterior, supongamos que conocemos que la cuota es diferente del 25%, entonces la formulación será como sigue:

$$H_0 : P = 25\% \text{ y la alternativa será } H_1 : P \neq 25\%$$

En Investigación Comercial se suele utilizar más el test de una cola, ya que se suele conocer el sentido de la afirmación. El bilateral se usa cuando no hay preferencia sobre la dirección del resultado.

7.3.2 ELECCIÓN DEL NIVEL DE SIGNIFICACIÓN

En el mundo científico, para que la hipótesis se convierta en teoría deberá ser validada mediante pruebas de hipótesis. Estas pruebas actúan como un filtro, eliminando las hipótesis falsas y manteniendo las verdaderas. Podemos afirmar que cuantas más pruebas haya superado una hipótesis, más convencimiento tendremos acerca de su veracidad.

Análisis de la Investigación Cuantitativa

En la praxis se ha escogido que el nivel de filtro sea del 95%, es decir, se eliminan el 95% de las hipótesis falsas.

En las pruebas de hipótesis se presentan cuatro situaciones que resumimos en la siguiente tabla:

	H_0 es verdadera	H_0 es falsa
No rechazar H_0	1. Decisión correcta Nivel de confianza Probabilidad $p = 1 - \alpha$	3. Error tipo II Probabilidad $p = \beta$
Rechazar H_0	2. Error tipo I Nivel de significación Probabilidad $p = \alpha$	4. Decisión correcta Poder de prueba Probabilidad $p = 1 - \beta$

1. La decisión correcta de no rechazar H_0 cuando es verdadera. El nivel de confianza proporciona el porcentaje de situaciones en los que la hipótesis nula se aceptaría siendo verdadera. No se refiere a la probabilidad de que los resultados sean observados en términos de muestreo, sino a la probabilidad de que la hipótesis sea verdadera con los datos obtenidos. La elección del nivel de significación y, por consiguiente, del nivel de confianza, es convencional y se fija a priori. En investigación comercial se suele trabajar con valores de nivel de significación $\alpha = 0.01$ (1%) y $\alpha = 0.05$ (5%). Los niveles de confianza vendrán dados por $1 - \alpha$. Es decir, fijamos a priori la probabilidad α de rechazar la H_0 siendo cierta. También en ocasiones se emplean niveles de significación menos exigentes, pudiendo llegar a veces hasta valores de α del 10%, que se denominan cuasi significativos. Rara vez se pasa de este valor del 10% para el nivel de significación.
2. Error de tipo I. Es el riesgo que se asume. El nivel de significación α representa el porcentaje de veces que se rechaza la H_0 cuando es verdadera.
3. Error de tipo II. No se rechaza H_0 cuando en realidad es falsa. A la probabilidad de que esto ocurra se le denomina β .

4. Se rechaza la hipótesis nula y es falsa. Se trata de una decisión correcta. La probabilidad de que esto ocurra se denomina poder o potencia de la prueba y viene dado por $1 - \beta$

La determinación de β es compleja. La relación entre α y β es de tipo inverso, de forma que al disminuir el error de tipo I aumenta β o la probabilidad de que ocurra un error tipo II, para una determinada muestra. El poder de la prueba $1 - \beta$ está ligada al valor del parámetro probado en la población objeto de estudio.

En la práctica, β se ignora o bien se determina después de estar seleccionada la muestra.

El interés del investigador se centra en no cometer un error de tipo I.

Como conclusión, podemos decir que en la práctica lo que se hace es fijar un nivel de significación (α) del 1% ó 5%. Y se rechaza la H_0 si hay menos del 1% ó 5% de posibilidades de que las diferencias obtenidas sean debidas al azar y se dice que la diferencia es significativa.

7.3.3 ELECCIÓN DEL TEST

La Estadística ha tenido un enorme desarrollo. En la actualidad, para cualquier diseño de investigación, nos encontramos con diversas pruebas estadísticas válidas para decidir acerca de H_0 .

Por ello, se hace necesario elegir criterios racionales para determinar qué prueba estadística es la más adecuada para analizar los datos de una Investigación Comercial. Se pueden utilizar diversos criterios. Los más usuales son:

la potencia, el procedimiento de obtención de los datos, el universo (N) de donde se obtuvo la muestra (n), las hipótesis que deseamos probar o la escala de medición utilizada.

VALIDEZ Y POTENCIA

La potencia de un análisis estadístico es, en parte, una función de la prueba estadística que se usa en un análisis. Una prueba estadística es válida si la probabilidad de rechazar H_0 , cuando ésta es verdadera, es igual al valor elegido α . Se dice que una prueba es potente si tiene gran probabilidad de rechazar H_0 cuando ésta es falsa.

MODELO ESTADÍSTICO

En el momento en el que identificamos la naturaleza del universo o población (N) y la técnica de muestreo a aplicar, hemos establecido un modelo estadístico. Cada prueba estadística se asocia a un modelo y a un requisito de medida.

Todas las decisiones tomadas para el uso de cualquier prueba estadística deben llevar consigo la siguiente condición.

“Si el modelo utilizado es correcto y los requisitos de medida fueron cumplidos, entonces....”.

Cuanto más pobres o débiles sean las suposiciones que definen el modelo, necesitaremos simplificar más la decisión alcanzada por la prueba estadística asociada al modelo y, por consiguiente, más generales serán las conclusiones.

Las pruebas más potentes son aquéllas que tienen las suposiciones más fuertes o extensas.

Las pruebas paramétricas (como la t, la Z o la F), tienen fuertes suposiciones subyacentes a su uso. Si estas suposiciones son válidas, las pruebas basadas en las mismas son las que tienen mayor probabilidad de rechazar H_0 cuando es falsa.

Hemos de resaltar que los requerimientos de los datos de investigación deben ser adecuados para la prueba.

EJEMPLO:

Las condiciones que debe satisfacer la prueba “t” para ser la más potente y aceptar así las conclusiones obtenidas con ella con el adecuado nivel de confianza, son:

1. Las observaciones deben ser independientes
2. Las observaciones deben derivar de poblaciones normalmente distribuidas
3. En el análisis de dos grupos, las poblaciones deben tener la misma varianza
4. Las variables serán medidas al menos en una escala de intervalo

Todas estas condiciones son elementos del modelo estadístico asociado con la distribución normal. De ordinario, estas suposiciones no son probadas en el curso del análisis estadístico, sino que son presunciones aceptadas y su certeza o falsedad determina la exactitud y significación de la probabilidad establecida mediante la prueba paramétrica.

Las pruebas paramétricas prueban hipótesis acerca de parámetros específicos. Se presupone que las hipótesis acerca de tales parámetros son idénticos a las hipótesis de investigación.

Cuando existen razones para creer que las condiciones se encuentran en los datos objeto de análisis, es posible elegir una prueba estadística paramétrica.

¿Qué ocurre si estas condiciones no se encuentran?

Cuando no se encuentran las suposiciones que constituyen el modelo estadístico de una prueba entonces esta no puede ser válida. Esto significa que el estadístico de prueba puede caer en la región de rechazo con una probabilidad mayor que alfa (α).

EFICACIA

Cuando comparamos dos muestras de tamaños diferentes, la potencia de la prueba aumenta conforme aumenta el tamaño de la muestra estadística. Por consiguiente, podremos utilizar una prueba menos potente con un tamaño de muestra más grande.

El concepto de potencia-eficacia se relaciona con el tamaño de la muestra que es necesario para lograr que la prueba B sea tan potente como la A cuando el nivel de significación y el tamaño de la muestra de la prueba A se mantienen constantes.

Si la prueba A es la prueba conocida más potente de su tipo y en prueba es una prueba para el mismo diseño de investigación, entonces

Potencia eficacia de prueba B

$$P_B = \frac{100n_a}{n_b}$$

Eficacia relativa asintótica de un estadístico

Es un modo de determinar el tamaño de muestra necesario para que la prueba B tenga la misma potencia que la prueba A.

Eficacia relativa

$$P_B = 100 \lim_{n_a \rightarrow \infty} \frac{n_a}{n_b}$$

ESCALA DE MEDIDA

Las escalas de medidas utilizadas en la Investigación Comercial son: nominal, ordinal, intervalo y de razón. Resmimos en la siguiente tabla sus características y relaciones:

Investigación Comercial

Escala	Relaciones que la definen
Nominal	Equivalencia
Ordinal	Equivalencia Mayor que
De intervalo	Equivalencia Mayor que Razón conocida entre cualesquiera intervalos
De razón	Equivalencia Mayor que Razón conocida entre cualesquiera intervalos Razón conocida entre cualesquiera de dos valores de la escala

Como conclusión, podemos decir que hemos de elegir el test más potente. Como regla general recordemos que:

- Los test unidireccionales son más potentes que los bidireccionales
- Los test paramétricos son más potentes que los no paramétricos
- Las pruebas que usan datos métricos son más potentes que las que no los utilizan

7.3.4 INTERPRETACIÓN DE LA PRUEBA

La prueba o test elegido tiene una variable asociada a un estadístico cuya distribución es conocida y cuyo valor se obtiene mediante la correspondiente fórmula con los datos de la muestra.

Para el nivel de significación α conocemos el valor teórico (valor de tablas) del test, que será el que compararemos con el obtenido de aplicar la fórmula a los datos obtenidos de la muestra.

Esta comparación es la que nos sirve como norma de decisión. En general

“Si el valor calculado es mayor que el teórico o de tablas para un determinado nivel de significación α , se rechaza la H_0 . En caso contrario no se rechaza la H_0 ”.

No rechazar la H_0 supone que los datos obtenidos en una investigación concreta son compatibles con la hipótesis propuesta. La interpretación ha de realizarse con prudencia

y criterio, recordando la frase de Henry Clay que dice: “las estadísticas no son sustitutos del criterio”.

En algunas pruebas, es preciso conocer los grados de libertad o el número de observaciones diferentes que obtenemos de una variable si descontamos el o los estadísticos calculados. En el caso de una tabla de r filas y c columnas, los grados de libertad son $(r - 1)$ para las filas, $(c - 1)$ para las columnas y $(r - 1)(c - 1)$ para la tabla.

RESUMEN

La metodología para realizar un contraste de hipótesis consiste en la realización de los siguientes pasos:

1. Definir H_0 y H_1
2. Definir un parámetro de medida que relacione los parámetros muestrales y poblacionales y la función de probabilidad
3. Establecer un criterio para juzgar si el valor calculado del parámetro es compatible con H_0
4. Realizar el muestreo
5. Obtener conclusiones de conformidad con lo expuesto

7.4 TIPOS DE TEST DE HIPÓTESIS

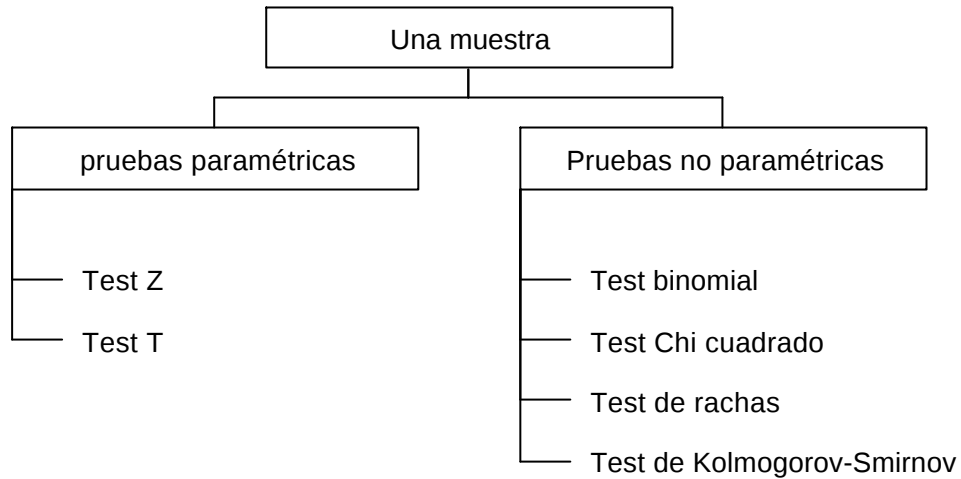
Como hemos visto ya, los test de hipótesis se utilizan para designar el procedimiento que contrastará la validez de una hipótesis. Son pruebas estadísticas que se utilizan para determinar si los resultados obtenidos con una muestra o dos elegidas al azar difieren marcadamente de aquellos que habría que esperar con la hipótesis planteada y la variación debida al muestreo.

Existen diferentes tipos de test. La forma más habitual de clasificarlos es la siguiente:

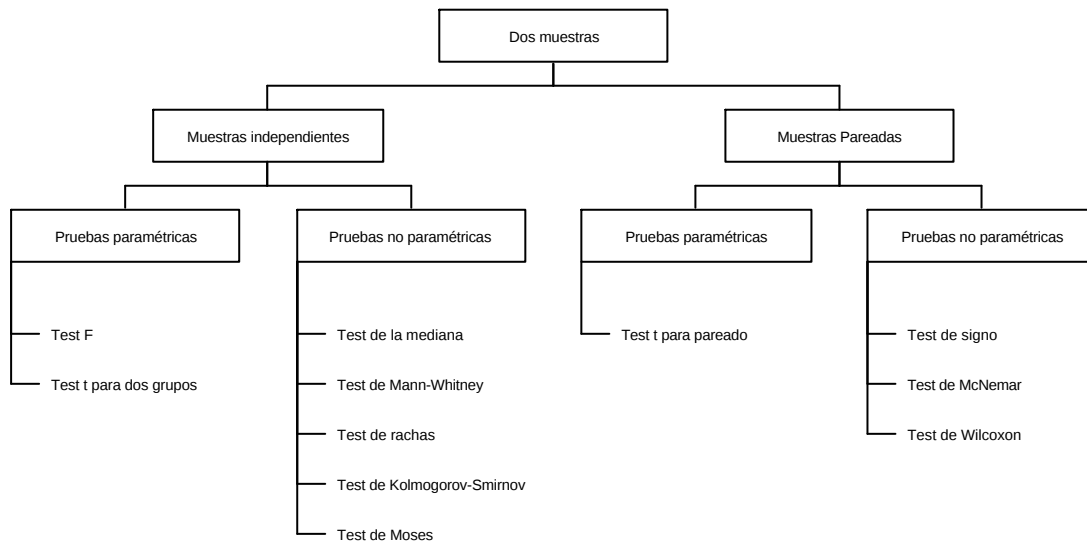
- Dependiendo del tipo de variables a contrastar, se distingue entre pruebas paramétricas y pruebas no paramétricas. Las primeras se aplican a variables medidas con escalas de ratio y de intervalo. Las segundas, a variables ordinales y nominales.
- Según realicemos el contraste sobre una muestra o sobre dos muestras, que a su vez se pueden diferenciar en muestras independientes o relacionadas (pareadas), aplicaremos diferentes tipos de test de hipótesis.

Los más usuales los resumimos en el siguiente esquema,
Investigación Comercial

Contraste de hipotesis



CONTRASTE DE HIPÓTESIS PARA DOS MUESTRAS



7.5 BREVE DESCRIPCIÓN DE LOS TEST

Vamos a reseñar brevemente en qué consiste cada uno de los procedimientos resumidos en el esquema anterior.

7.5.1 CONTRASTES PARA UNA MUESTRA

1. Pruebas paramétricas

Test Z. Está basado en la distribución normal. Se utiliza para contrastar estadísticos de una distribución con respecto a valores de parámetros del universo. El paquete estadístico calcula el valor de Z y luego calcula la probabilidad asociada al mismo. Se compara la probabilidad con el nivel de significación seleccionado y no se rechaza la H_0 en el caso de que la probabilidad sea mayor que dicho nivel. Se acepta la hipótesis alternativa cuando la probabilidad es menor que el nivel de confianza.

Test t. Se basa en la t de Student. Parte del conocimiento de la media de la muestra y de la desviación típica de la media y de la media del universo. La interpretación de los resultados es idéntica al test Z.

2. Pruebas no paramétricas

Test binomial. Se utiliza para variables dicotómicas, con la finalidad de contrastar si una variable procede de una población binomial con una probabilidad determinada de que se produzca un suceso.

Test Chi cuadrado. Se suele utilizar en variables medidas en escala nominal. Se basa en la distribución Chi cuadrado.

Test de rachas (test runs). Las rachas se definen como el número de veces que se produce el cambio de un valor a otro dentro de la distribución de la variable dicotómica. Su finalidad es comprobar si la variable se distribuye aleatoriamente o no. Se basa en el orden de ocurrencia de los dos valores de una variable de tipo dicotómico.

Test de Kolmogorov-Smirnov. Se utiliza para contrastar la hipótesis de que la muestra estudiada se distribuye como alguna de las principales distribuciones, normal, uniforme o de Poisson.

7.5.2 CONTRASTES PARA DOS MUESTRAS INDEPENDIENTES

1. Pruebas paramétricas

Test t para dos grupos independientes. Se aplica a variables medidas en escalas de intervalo o de razón. Se utiliza para contrastar si dos muestras independientes proceden de poblaciones con la misma media.

Test F (F de Barlett Box). Se utiliza para contrastar si las muestras proceden de poblaciones con idénticas varianzas. Realiza el cociente entre las medias cuadráticas de las dos muestras, teniendo en cuenta el número de grados de libertad del numerador y del denominador, asignando así la probabilidad.

2. Pruebas no paramétricas

Test de la mediana. Se usa para contrastar si dos o más muestras independientes pertenecen a poblaciones con la misma mediana.

Test de rachas de Wald-Wolfowitz. Contrasta la hipótesis de que dos muestras proceden de la misma población. Se aplica normalmente a variables medidas en escala ordinal. Cuando existe un número elevado de rachas, las muestras pertenecen a la misma población.

Test de Kolmorov-Smirnov (para dos muestras). Sirve para contrastar si las dos muestras provienen de la misma distribución. El test utiliza las diferencias existentes entre mediana, dispersión, asimetría e incluso más estadísticos de ambas muestras, comparando las distribuciones acumuladas y las diferencias con patrones conocidos.

Test de Mann-Whitney. Se aplica a variables medidas en escala ordinal. Es una prueba parecida a la t.

Test de Moses. Se aplica a variables medidas en escala ordinal. Consiste en contrastar la hipótesis de que la variable estudiada afecta a unos elementos en una dirección y a otros en la dirección opuesta.

7.5.3 CONTRASTES PARA DOS MUESTRAS RELACIONADAS

1. Pruebas paramétricas

Test t para pareado. Se utiliza para comprobar si dos muestras provienen de poblaciones con igual media.

2. Pruebas no paramétricas

Test de McNemar. Se aplica sobre dos variables dicotómicas relacionadas. El contraste lo realiza a través de la distribución Chi cuadrado, analizando la probabilidad de ocurrencia de las situaciones posibles (0,1) y (1, 0).

Test del signo. Se utiliza para contrastar la hipótesis de que las dos variables tienen idéntica distribución. Si las dos muestras tienen la misma distribución, la mitad de las diferencias deberían ser positivas y la otra mitad negativas; por consiguiente el método que sigue se basa en la dirección (signo) de las diferencias entre las dos variables.

Test de Wilcoxon. Se utiliza para contrastar la hipótesis de que las dos variables tienen idéntica distribución. Tiene en cuenta la magnitud de las diferencias dentro de los pares y pondera con un valor mayor a los pares que presentan mayores diferencias.

8. ANÁLISIS BIVARIABLE

El análisis bivariable trata de analizar la relación entre dos variables. Este análisis permite comprobar si existe asociación entre esas variables así como medir la fuerza de esa asociación. Esta asociación no implica necesariamente causalidad (no es preciso que se cumpla que la variable causa preceda a la variable efecto).

Como las variables a estudiar pueden ser de tipo cuantitativo y cualitativo, las situaciones y principales técnicas de estudio se recogen en la tabla siguiente

TIPO DE VARIABLE	CUANTITATIVA	CUALITATIVA
CUANTITATIVA	Coeficiente de correlación (Pearson) Regresión Coeficiente alfa (a) de Cronbach	Análisis de varianza (Test F) Test “t” de medias
CUALITATIVA		Tabla de contingencia (Chi cuadrado) Correlación de rangos (Spearman) Coeficiente de Cramer

8.1 RELACIÓN ENTRE DOS VARIABLES CUALITATIVAS

8.1.1 TABLAS DE CONTINGENCIA

CONCEPTO

La forma más frecuente de presentación del análisis es la tabulación cruzada, en ella se presentan las dos variables cualitativas con un número no excesivamente grande de categorías, en la praxis 4 ó 5 como mucho.

Una tabla de contingencia (Crosstab) es el resultado de clasificar los elementos de la muestra con arreglo a dos variables cualitativas (nominales), cada una de ellas diversificadas en modalidades o categorías mutuamente excluyentes.

Con la tabla de contingencia se puede verificar si existe relación de dependencia entre las dos variables. En el punto siguiente indicamos un ejemplo de una tabla de contingencia obtenida con el paquete estadístico SPSS para Windows

Análisis de la Investigación Cuantitativa

8.1.1.1 TABLA DE CONTINGENCIA : SEXO Y CONOCIMIENTO DE INFORMÁTICA

Las variables que se cruzan son sexo de los encuestados y conocimiento de informática de los mismos. La tabla resultante del tratamiento estadístico es la siguiente:

Tabla:Sexo by conoci.

Count	Ninguno	Poco	Bueno	Elevado	Experto	
Exp Val						
Row Pct						Row
Col Pct						Total
Tot Pct	1	2	3	4	5	
Mujer	9	9	12	3	3	36
	7'1	6'2	6'7	12'0	4'0	44'4%
1	25'0%	25'0%	33'3%	8'3%	8'3%	
	56'3%	64'3%	80'0%	11'1%	33'3%	
	11'1%	11'1%	14'8%	3'7%	3'7%	
Hombre	7	5	3	24	6	45
	8'9	7'8	8'3	15'0	5'0	55'6%
2	15'6%	11'1%	6'7%	53'3%	13'3%	
	43'8%	35'7%	20'0%	88'9%	66'7%	
	8'6%	6'2%	3'7%	29'6%	7'4%	
Column	16	14	15	27	9	81
Total	19'8%	17'3%	18'5%	33'3%	11'1%	100'0%

Chi-square	Value	DF	Significance
Pearson	23'41527	4	'00010
Likelihood Ratio	25'80228	4	'00003
Mantel-Haenszel test for linear association	8'74396	1	'00311

Minimum Expected Frequency 4'000

Cells with Expected Frequency < 5 1 OF 10 (10'0%)

Number of Missing Observations: 21

8.1.1.2 SIGNIFICADO DE LOS ELEMENTOS QUE COMPONEN LA TABLA

Las dos variables cruzadas en esta tabla, sexo (A) y conocimiento de la informática (B), se dividen en las siguientes categorías:

Variable A, 2 categorías: 1 (Mujer) y 2 (Hombre).

Variable B, 5 categorías: 1 (Ninguno), 2 (Poco), 3 (Bueno), 4 (Elevado), y 5 (Experto).

Las cifras enmarcadas indican el número de observaciones totales válidas. En este caso, hay 81 casos, lo que representa el 100%.

En la columna titulada *Row Total* (Total fila) aparecen el número de observaciones de cada una de las categorías de la variable A, tanto en valores absolutos (frecuencias absolutas) como en porcentajes sobre el total casos (frecuencias relativas o probabilidad). Así, en el ejemplo anterior, el número de mujeres que han contestado la encuesta es de 36, que representa el 44,4% del total de casos válidos. El número de hombres es de 45 y su porcentaje el 55,6%.

Evidentemente, la suma de las frecuencias absolutas de cada categoría de esta variable coincide con el total de casos ($36 + 45 = 81$), y las frecuencias relativas suman 100.

En la fila titulada *Column Total* (Totales de la columna) tenemos el total de casos que pertenecen a cada una de las categorías de la variable B. Así, en el ejemplo, el total de la primera columna nos indica que 16 personas (hombres+mujeres) no tienen ningún conocimiento de la informática (frecuencia absoluta). Esto representa el 19,8% de las observaciones válidas ($16 : 81$) % (frecuencia relativa o probabilidad).

Como en el total de la fila, la suma de las frecuencias absolutas de cada categoría de esta variable coincide con el total de casos ($16 + 14 + 15 + 27 + 9 = 81$), y las frecuencias relativas suman 100.

Dentro de cada casilla o celda obtenemos datos correspondientes a la intersección (cruce) entre las categorías de las variables (A_i y B_j). La información contenida en cada casilla queda reflejada en el margen superior izquierdo de la tabla (no tiene por qué ser siempre la misma). En este ejemplo, las cifras de cada celda corresponden, correlativamente, a:

Count (Frecuencia) es el número de casos que han contestado a la vez, a la categoría i de la variable A y la categoría j de la variable B. En el ejemplo 9 es el número de mujeres que han contestado “Ningún conocimiento de informática”.

La suma de las frecuencias de A en una modalidad y que cumplen las diferentes categorías de B, nos da la frecuencia de A en esa modalidad. En el ejemplo, para la categoría mujer, sería: $9 + 9 + 12 + 3 + 3 = 36$

Exp Val (Valor esperado.) Corresponde al número de casos que deberían aparecer en la casilla, si las dos variables fueran independientes entre sí. Luego se verá cómo se calcula y su utilidad.

Row Pct (Porcentaje fila). Porcentaje de casos de un cruce sobre el total de casos de la fila (observaciones de la categoría i de la variable A). Las 9 observaciones del cruce Mujer y Ningún conocimiento, son un 25% del total de Mujeres. O dicho de otra forma, del total de mujeres, un 25% tiene “Ningún conocimiento de informática”. Por tanto, esta frecuencia relativa es una probabilidad condicionada: probabilidad de B_j condicionada a A_i , $P(B_j / A_i)$.

Col Pct (Porcentaje columna). Porcentaje de casos de un cruce sobre el total de casos de la columna (observaciones de la categoría j de la variable B). En el ejemplo, las 9 observaciones del cruce Mujer y Ningún conocimiento, son un 56,3% del total de personas que tienen “Ningún conocimiento de informática”. O bien, un 56,3% de los que tienen “Ningún conocimiento de informática” son mujeres. Esta frecuencia relativa también es una probabilidad condicionada. En este caso, la probabilidad de A_i condicionada a B_j , $P(A_i / B_j)$.

Tot Pct (Porcentaje total). Porcentaje de casos de un cruce sobre el total de casos. Por tanto, equivale a la frecuencia relativa o probabilidad de la intersección de A_i con B_j , $P(A_i \cap B_j)$. En el ejemplo, las 9 observaciones del cruce Mujer y Ningún conocimiento, son un 11,1% del total de casos válidos ($9 : 81$)%.

8.1.1.3 ANÁLISIS DE INDEPENDENCIA ENTRE LAS DOS VARIABLES. ESTADÍSTICO χ^2 (CHI CUADRADO)

El análisis de independencia entre las dos variables se hace con el estadístico χ^2 (Chi cuadrado) que aparece después de la tabla. De las tres líneas que proporciona el ordenador hay que fijarse únicamente en la primera, la χ^2 de Pearson, que es la más genérica.

Para calcular la χ^2 primero hay que determinar el Valor esperado de una celda. Tal como hemos indicado anteriormente, el Valor esperado es el número de casos que deberían aparecer en una casilla si las dos variables fueran independientes.

El Valor esperado es igual al número de observaciones del total de la fila a la que pertenece la casilla, multiplicado por el número de observaciones del total de la columna y dividido por el número de observaciones válidas. En nuestro ejemplo, para la primera celda, el Valor esperado es igual a 7,1 obtenido con la siguiente operación matemática: $(36 \times 16) / 81 = 7,1$.

También podemos obtener este valor si aplicamos el porcentaje del total de la fila al número de observaciones total de la columna ($44,4 \% \times 16 = 7,1$) o el porcentaje del total de la columna al número de observaciones total de la fila ($19,8 \% \times 36 = 7,1$). Con estas nuevas fórmulas se puede interpretar el Valor esperado de una manera más simple. Si las dos variables son independientes el número de mujeres con “Ningún conocimiento de informática” debería ser igual al porcentaje de personas que tienen “Ningún conocimiento de informática” (19,8%) multiplicado por el número de mujeres de la muestra (36) o al porcentaje de mujeres (44,4%) aplicado al total de personas que tienen “Ningún conocimiento de informática” (16).

De forma genérica, si las dos variables son independientes, los porcentajes del total de la fila (*Row Total*) deberían ser iguales en cada columna en los porcentajes sobre el total de la columna (*Col Pct*), ya que la variable representada en las filas no influye en la otra variable. De igual forma, los porcentajes del total columna (*Column Total*) deberían ser los mismo en cada fila que los porcentajes sobre el total de la fila (*Row Pct*). En nuestro ejemplo, el 44,4% de mujeres y el 55,6% de hombres debería repetirse para todas las categorías de la variable “Conocimiento de informática”. O bien, los porcentajes de “Conocimiento de informática” correspondientes a las personas (Total columna), deben aparecer tanto para hombres como para mujeres.

La columna del ejemplo en la situación teórica de independencia entre las dos variables quedaría como sigue:

Mujeres 1	Count	7'1
	Exp Val	7'1
	Row Pct	19'8
	Col Pct	44'4
	Tot Pct	8'8
Hombres 2	Count	8'9
	Exp Val	8'9
	Row Pct	19'8
	Col Pct	55'6
	Tot Pct	11'0

La χ^2 se calcula de la siguiente manera:

$$\sum_{i=1}^r \sum_{j=1}^c \frac{[n_{ij} - E(n_{ij})]^2}{E(n_{ij})}$$

donde los sumatorios recogen los valores de todas las celdas, siendo n_{ij} la frecuencia observada (*Count*) y $E(n_{ij})$ el valor esperado (*Exp Val*) de cada celda.

La χ^2 no puede utilizarse si cualquiera de las celdas tiene un valor esperado menor que 1 o si más de un 20% de las celdas tienen un valor esperado menor que 5. En nuestro ejemplo, podemos aplicarla, ya que el valor mínimo esperado es 4 y sólo en un 10% de las celdas (1 de 10) hay un valor esperado menor que 5.

Minimum Expected Frequency - 4.000

Cells with Expected Frequency < 5 - 1 OF 10 (10.0%)

La última línea del resultado de una tabla de contingencia nos indica el número de observaciones “missing”, es decir, aquéllas para las que el encuestado no ha contestado a alguna de las variables cruzadas o a ambas. En el ejemplo, en 21 casos no se dispone de información referida a las variables cruzadas.

8.1.2 COEFICIENTE V DE CRAMER

Se trata de un coeficiente que toma el valor 1 cuando hay una asociación perfecta entre los atributos considerados, con independencia del número de filas y columnas de la tabla de contingencia.

Su fórmula matemática es:

$$V = \sqrt{\frac{c^2}{m n}}$$

Donde m es el mínimo de los grados libertad de las filas y columnas. Por tanto

$$m = \min. (h - 1, k - 1)$$

n es el tamaño de la observación (muestra) y χ^2 es el estadístico Chi cuadrado.

8.1.3 CORRELACIÓN DE RANGOS DE SPEARMAN

Este coeficiente determina la correlación entre dos variables cualitativas medidas en escala ordinal. Este coeficiente varía entre -1 y $+1$; un valor positivo nos indica una correlación en el mismo sentido (directa), mientras que los valores negativos ponen de manifiesto una relación inversa.

Para su determinación se aplica la fórmula de Spearman:

$$r = 1 - \frac{6 \sum_{i=1}^n D_i^2}{N(N^2 - 1)}$$

Donde N es el número de pares de observaciones y D la diferencia entre los grados de los valores correspondientes (x e y).

8.2 MÉTODOS DE MEDICIÓN EN EL ANÁLISIS ENTRE DOS VARIABLES CUANTITATIVAS

La asociación entre dos variables no implica necesariamente una relación causal, sino simplemente conociendo la asociación entre dos variables se podrá anticipar la variación de una variable conociendo el comportamiento de la otra.

A continuación vamos a reseñar las medidas de asociación más habituales en la Investigación Comercial.

8.2.1 CORRELACIÓN

Es la asociación entre las variaciones de los valores de dos variables. La asociación puede ser directa (mismo sentido) o inversa (sentidos opuestos).

La asociación es la relación entre el comportamiento de las dos variables. Esta asociación puede deberse a la “casualidad” o bien ser causal.

Se dice que la relación es causal cuando un cambio en una de las variables (independiente o explicativa) produce un cambio en la otra (dependiente).

Para que exista causalidad se deben cumplir al menos las siguientes condiciones⁷:

- Variación Concomitante. Supone que ambas variables varían a la vez.
- Temporalidad en la variación. La causa debe ir por delante del efecto.
- Control sobre otros factores distorsionadores. Se deben encontrar las posibles correlaciones espurias (correlaciones debidas a la asociación entre dos variables, bien de forma accidental o casual, bien por efecto de una tercera variable, sin que se produzca ninguna relación causal entre ellas).

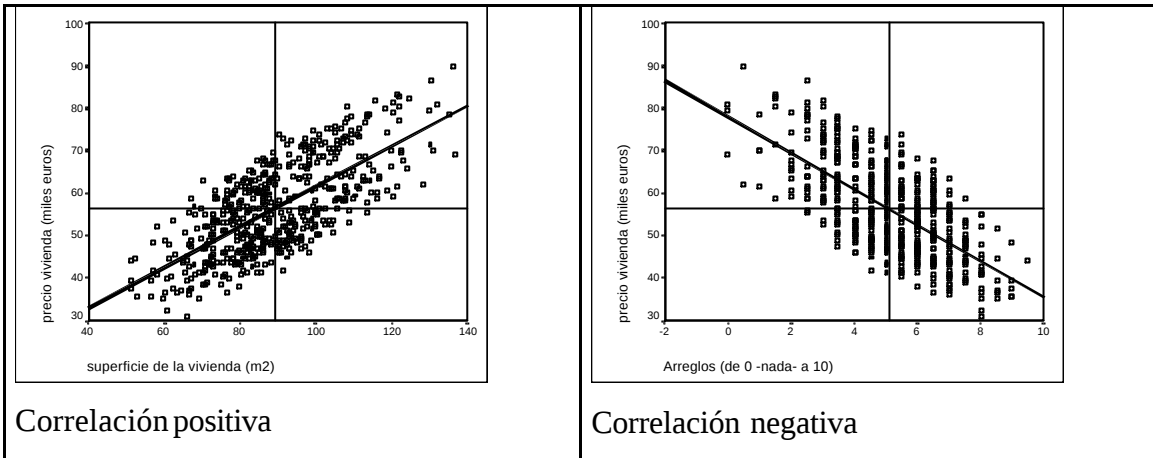
La correlación entre dos variables cuantitativas se estudia habitualmente mediante el análisis de correlación lineal, que se mide mediante el coeficiente de Pearson.

Este coeficiente puede tomar valores entre +1 y -1.

El valor 0 indica la ausencia de correlación. Si el coeficiente es positivo indica que la correlación es directa (en el mismo sentido), en caso contrario se dice que es indirecta.

⁷ Dillon y otros Marketing Research in a Marketing environment Irwin 1994
Investigación Comercial

Gráfica



El coeficiente entre las dos variables x e y se calcula dividiendo la covarianza de ambas variables entre el producto de las desviaciones típicas de ambas variables.

$$r = \frac{s(x,y)}{s(x)s(y)}$$

$$r = \frac{n \sum xy - \sum x \sum y}{\sqrt{(n \sum x^2 - (\sum x)^2)(n \sum y^2 - (\sum y)^2)}}$$

Análisis de la Investigación Cuantitativa

8.2.1.1 EJEMPLO

Supongamos que queremos conocer la correlación existente entre la inversión en promoción de ventas (x) y la cifra de ventas (y), en miles de euros

Los datos los resumimos en la siguiente tabla

X	y	xy	x ²	y ²
2	50	100	4	2500
2'5	70	175	6'25	4900
3	75	225	9	5625
3	80	240	9	6400
4	85	340	16	7225
3'5	90	315	12'25	8100
Σx = 18	Σy = 450	Σxy = 1395	Σx ² = 56'5	Σy ² = 34750

El coeficiente de correlación lineal (s e u o) será

$$r = \frac{6 \bullet 1395 - 18 \bullet 450}{\sqrt{(6 \bullet 56'6 - 18^2)(6 \bullet 34750 - 450^2)}} = 0'88$$

Por tanto, podemos suponer que existe muy buena relación entre ambas variables.

Para concluir si la r de Pearson es significativa al 5% ó al 1%, deberemos consultar con la correspondiente tabla.

Análisis de la Investigación Cuantitativa

8.2.1.2 TABLA COEFICIENTE DE CORRELACIÓN R DE PEARSON

Prueba unilateral: nivel de significación				
	0,05	0,025	0,01	0,005
Prueba bilateral: nivel de significación .				
Grados de libertad	0,10	0,05	0,02	0,01
1	0,988	0,997	0,999	0,999
2	0,900	0,195	0,980	0,990
3	0,805	0,878	0,934	0,959
4	0,729	0,811	0,882	0,917
5	0,669	0,784	0,833	0,874
6	0 622	0,707	0,789	0,834
7	0,582	0,666	0,750	0,798
8	0,549	0,632	0,716	0,765
9	0,521	0,602	0,685	0,735
10	0,497	0,576	0,658	0,708
11	0,576	0,553	0,634	0,684
12	0,458	0,532	0,612	0,661
13	0,441	0,514	0,592	0,641
14	0,426	0,497	0,574	0,623
15	0,412	0,482	0,558	0,606
16	0,400	0,468	0,542	0,590
17	0,389	0,456	0,528	0,575
18	0,378	0,444	0,516	0,561
19	0,369	0,433	0,503	0,549
20	0,360	0,423	0,492	0,537
21	0,352	0,413	0,482	0,526
22	0,344	0,404	0,472	0,515
23	0,337	0,396	0,462	0,505
24	0,330	0,388	0,453	0,496
25	0,323	0,381	0,445	0,487
26	0,317	0,374	0,437	0,479
27	0,311	0,367	0,430	0,471
28	0,306	0,361	0,423	0,463
29	0,301	0,355	0,416	0,486
30	0,296	0,349	0,409	0,449
35	0,275	0,325	0,381	0,418
40	0,257	0,304	0,358	0,393
45	0,243	0,288	0,338	0,372
50	0,231	0,273	0,322	0,354
60	0,211	0,250	0,295	0,325
70	0,195	0,232	0,274	Q303
80	0,183	0,217	0,256	0,283
90	0,173	0,205	0,242	0,267
100	0,164	0,195	0,230	0,254

8.2.2 REGRESIÓN SIMPLE

En términos generales, podemos definir la regresión lineal simple como el estudio entre una variable a explicar con respecto a otra, que denominamos explicativa. Ambas variables deben ser cuantitativas.

El modelo de regresión lineal de primer orden viene expresado por la siguiente ecuación:

$$y = \beta_0 + \beta_1 x + \varepsilon$$

donde:

y = variable **dependiente** o variable a explicar

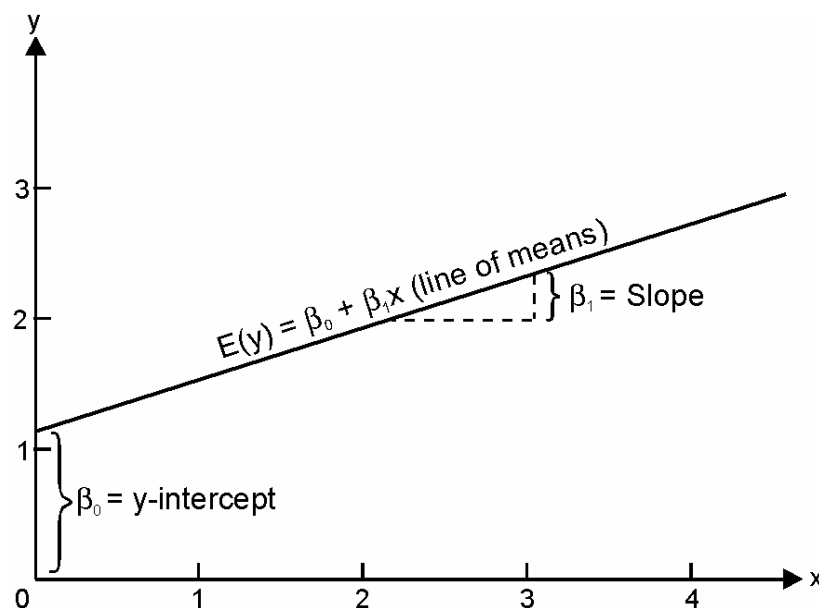
x = variable **independiente** o variable explicativa

ε (épsilon) = **error** o perturbación aleatoria

β_0 = **origen** de la recta: punto donde la recta corta el eje de ordenadas o eje de la y

β_1 = **pendiente** de la recta o coeficiente de regresión: nos indica cuánto aumenta o disminuye la variable dependiente por cada incremento en 1 unidad de la variable independiente

La correspondiente representación gráfica es:



8.2.2.1 OBJETIVOS

Los objetivos que se pretenden con este tipo de análisis son varios:

Determinar la forma de la relación entre las dos variables, dependiente e independiente

Comprobar la hipótesis

Predecir los valores que tomará la variable dependiente

Para dar satisfacción a estos objetivos nos interesa saber:

Cómo se calculan los coeficientes de regresión, β_0 y β_1

Cómo se interpretan

Cómo se determina si son o no estadísticamente significativos

Cómo se comprueban las hipótesis del modelo

Estimación del modelo de regresión por mínimos cuadrados ordinarios

Con los datos de la muestra se pueden estimar los parámetros desconocidos del modelo del siguiente modo:

$$y = \beta_0 + \beta_1 x + \varepsilon$$

$$\hat{y} = \hat{\beta}_0 + \hat{\beta}_1 x$$

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_i$$

$$y_i - \hat{y}_i = y_i - (\hat{\beta}_0 + \hat{\beta}_1 x_i)$$

$$\text{SSE (Suma de Errores al Cuadrado)} = \sum_{i=1}^n [y_i - (\hat{\beta}_0 + \hat{\beta}_1 x_i)]^2$$

La recta de mínimos cuadrados ordinarios es, precisamente, aquélla que minimiza la suma de los errores cuadrados.

Fórmulas para obtener los estimadores mínimos cuadrados

$$\text{Pendiente: } \hat{\beta}_1 = \frac{SS_{xy}}{SS_{xx}} \quad \text{Origen: } \hat{\beta}_0 = \bar{y} - \hat{\beta}_1 \bar{x}$$

donde

$$SS_{xy} = \sum_{i=1}^n x_i y_i - \frac{\left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n y_i \right)}{n}$$

$$SS_{xx} = \sum_{i=1}^n x_i^2 - \frac{\left(\sum_{i=1}^n x_i \right)^2}{n}$$

n = tamaño de la muestra

Los estimadores mínimos cuadrados son los mejores que se pueden conseguir (insesgados, eficientes y consistentes) si se cumplen ciertas hipótesis sobre ε (las perturbaciones).

8.2.2.2 EJEMPLO REGRESIÓN LINEAL

En los últimos diez años, las ventas de un fabricante respecto de un determinado producto han sido las siguientes:

Año	1	2	3	4	5	6	7	8	9	10
Ventas	40	42	50	49	47	50	52	51	55	60

Determinar la recta de regresión (tendencia).

La recta vendrá dada por:

$y = a + b x$ donde y son las ventas x los años

Las ecuaciones a resolver son :

$$\sum y_i = na + b \sum x_i$$

$$\sum x_i y_i = a \sum x_i + b \sum x_i^2$$

Análisis de la Investigación Cuantitativa

Para resolver el correspondiente sistema de ecuaciones realizamos la siguiente tabla

X	y	xy	x ²
1	40	40	1
2	42	84	4
3	50	150	9
4	49	196	16
5	47	235	25
6	50	300	36
7	52	364	49
8	51	408	64
9	55	495	81
10	60	600	100
55	496	2872	385

El sistema de ecuaciones, a resolver queda como sigue:

$$496 = 10 \alpha + \beta 55$$

$$2872 = 55 \alpha + \beta 385$$

Luego

$$2872 = 55 \left(\frac{496 - 55b}{10} \right) + 385b$$

Una vez realizados los correspondientes cálculos obtenemos que $\beta = 1'745$ y $\alpha = 40$

La recta de regresión será: $y = 40 + 1'745 x$

8.2.3 COEFICIENTE ALFA DE CRONBACH

El coeficiente alfa de Cronbach es un estimador de la consistencia interna de una escala de medida. La expresión matemática de este coeficiente es:

$$a = \frac{k}{k-1} \left(1 - \frac{\sum_{i=1}^k s_i^2}{s_s^2} \right)$$

Donde k es el número de ítem de la escala, s_i^2 es la varianza del ítem “i” y s_s^2 es la varianza de toda la escala.

El valor de alfa tiende a aumentar a medida que se incrementa el número de ítem de la escala. Este coeficiente varía entre 0 y 1, y puede tomar valores negativos cuando existen ítem correlacionados negativamente, en esta situación el coeficiente alfa de Cronbach no es el adecuado para medir la confiabilidad (grado en el que la medida de una variable está libre de error aleatorio y por consiguiente proporciona resultados consistentes).

La principal aplicación de este índice es medir la confiabilidad de una escala.

Un valor del coeficiente alfa de Cronbach inferior a 0.7 indica una baja consistencia interna.

Cuando $\alpha < 0.7$ nos indica que la escala mide varios fenómenos y, por consiguiente, puede no ser apropiada para los objetivos de la investigación; en estas circunstancias de escalas multidimensionales suele ser mejor aplicar análisis factorial.

8.3 RELACIÓN ENTRE UNA VARIABLE CUANTITATIVA Y OTRA CUALITATIVA

8.3.1 ANÁLISIS DE LA VARIANZA

El análisis de la varianza, en el análisis bivariable, es una técnica estadística que relaciona una variable cualitativa con otra cuantitativa.

Determina la existencia de diferencias significativas entre las medias de una variable dependiente.

Es de gran utilidad en la investigación experimental.

La variable dependiente debe estar medida en escala métrica, mientras que la variable independiente es cualitativa. El proceso es como sigue:

El análisis de la varianza (ANOVA) se utiliza cuando no conocemos previamente el comportamiento de la variable dependiente (cuantitativa) sin la influencia del factor principal controlado (variable cualitativa).

8.3.1.1 MÉTODO ANOVA TRADICIONAL UNIDIRECCIONAL

En este método, el test de significación que se utiliza es el F, que compara los efectos de los diferentes tratamientos recogidos por la varianza de la dispersión factorial con los efectos de los factores controlados representados por la varianza de la dispersión residual a través de una relación.

$$F = \frac{\text{Varianza.de.la.dispersión.factorial}}{\text{Varianza.de.la.dispersión.residual}}$$

El examen de las diferencias entre las medias implica la descomposición de la variación total observada en la variable dependiente. Luego será:

$$\text{Dispersión total} = \text{Dispersión factorial} + \text{Dispersión residual}$$

Y las varianzas de las dispersiones son:

1- Dispersión total

Viene dada por la siguiente fórmula:

$$S_t^2 = \frac{\sum_{ij} (x_{ij} - m)^2}{n - 1}$$

donde n es el tamaño de la muestra.

x_{ij} es el valor de cada uno de los datos individuales considerados

m es el valor de la media general

2- Dispersión factorial

Su expresión matemática es:

$$S_f^2 = \frac{\sum_i (m_i - m)^2}{K - 1}$$

donde m_i es la media de los diferentes grupos estudiados,

m es la media general y

k es el número de grupos considerados

3- Dispersión Residual

Su formulación es:

$$S_r^2 = \frac{\sum_{ij} (x_{ij} - m_i)^2}{n - K}$$

donde k es el número de grupos considerados

m_i es la media de los diferentes grupos estudiados

x_{ij} es el valor de cada uno de los datos individuales considerados

n es el tamaño de la muestra

Cuando el valor de F es igual a la unidad o muy próximo a este valor, significa que las dos dispersiones apenas difieren y, como consecuencia, se puede afirmar que los diferentes tratamientos no han tenido eficacia.

Si el valor de F es muy superior a la unidad, se admite que el efecto de los tratamientos ha sido eficaz.

Como el número de observaciones que se realizan es limitado, no podemos conocer el valor exacto de la varianza. Por ello, el valor de F oscilará en torno a la unidad como consecuencia de las variaciones del muestreo.

Este efecto se mitiga con la utilización de las tablas de Snedecor. De esta manera, cuando el valor de F calculado es superior al valor crítico de F indicado en las tablas, significa que los tratamientos aplicados son eficaces.

Si el valor calculado de F es inferior al valor obtenido en las tablas, esta diferencia se debe a las variaciones de muestreo.

Análisis de la Investigación Cuantitativa

TABLA ESTADÍSTICA: DISTRIBUCIÓN DE LA F (nivel de confianza 95%)

“n”	“m”				
	1	2	3	4	5
1	161´4	199´5	215´7	224´6	230´2
2	18´51	19	19´16	19´25	19´30
3	10´13	9´55	9´28	9´12	9´01
4	7´71	6´94	6´59	6´39	6´26
5	6´61	5´79	5´41	5´19	5´05
6	5´99	5´14	4´76	4,53	4´39
7	5´59	4´74	4´35	4´12	3´97
8	5´32	4´46	4´07	3´84	3´69
9	5´12	4´26	3´86	3´63	3´48
10	4´96	4´10	3´71	3´48	3´33
11	4´84	3´98	3´59	3´36	3´20
12	4´75	3´89	3´49	3´26	3´11
13	4´67	3´81	3´41	3´18	3´03
14	4´6	3´74	3´34	3´11	2´96
15	4´54	3´68	3´29	3´06	2´90

Siendo **m** los grados de libertad del numerador y **n** los grados de libertad del denominador.

8.3.1.1.1 EL PROCESO DEL MÉTODO ANOVA UNIDIRECCIONAL

Se determinan las siguientes dispersiones:

1.- Dispersión total (DT)

Mide la suma de las dispersiones.

2.- Dispersión factorial (DF)

Mide la dispersión entre los grupos creados por las diferentes alternativas del factor o factores estudiados.

Dependiendo del tipo de experimento, pueden existir varias dispersiones factoriales, correspondientes al factor principal y a los factores de bloque.

3.- Dispersión residual (DR)

Mide la dispersión dentro de los grupos creados por las diferentes alternativas del factor o factores estudiados.

$$DT = DF + DR \quad DR = DT - DF$$

4.- Se calcula el cuadrado medio total (CMT)

Se trata de la dispersión total dividida por el número de grados de libertad

$$CMT = DT / gl \quad \text{donde } gl \text{ son los grados de libertad.}$$

5.- Se calcula el cuadrado medio factorial (CMF)

Se trata de la dispersión factorial dividida por el número de grados de libertad.

$$CMF = DF / gl$$

Dependiendo del tipo de experimento pueden existir varias varianzas factoriales, que corresponden al factor principal y a los factores bloque.

6.- Se calcula el cuadrado medio residual (CMR)

Se trata de la dispersión residual dividida por el número de grados de libertad.

$$CMR = DR / gl$$

7.- Se realiza el test de la F.

Para cada factor estudiado se calcula:

7-1.- Se calcula el estadístico F.

$$F = CMF / CMR$$

Si el valor de F es menor que uno, es decir $CMF < CMR$, no existe un efecto significativo del factor estudiado sobre la variable dependiente, y por tanto no es necesario realizar la comparación de F con el correspondiente valor de las tablas.

Análisis de la Investigación Cuantitativa

7-2.- Se determina el valor de F en las tablas estadísticas de la distribución de la F, en base a los grados de libertad del numerador y del denominador.

7-3.- Se comparan ambos valores.

La hipótesis nula H_0 es: NO EXISTE EFECTO SIGNIFICATIVO DEL FACTOR ESTUDIADO.

Entonces:

Si $F > F_t$ (tabla), se rechaza H_0 y por tanto el factor estudiado tiene una influencia significativa sobre la variable dependiente.

Si $F = F_t$ (tabla), no se rechaza H_0

EJEMPLO DE ANOVA UNIDIRECCIONAL

Vamos a desarrollar lo expuesto anteriormente mediante un caso práctico:

Un banco lleva a cabo un experimento comercial realizando tres promociones diferentes para el lanzamiento de un nuevo producto. Estas promociones consisten en:

P1: regalo de una bicicleta, P2: regalo de un ordenador, P3: regalo de los electrodomésticos de la cocina.

Cada promoción se prueba en cinco sucursales diferentes durante un mes. Los resultados obtenidos, en cuanto a unidades de producto colocadas entre la clientela, se recogen en el cuadro siguiente:

	S1	S2	S3	S4	S5
P1	65	50	30	40	65
P2	30	25	15	20	35
P3	15	10	10	25	50

SOLUCIÓN

Definiremos las siguientes características:

Factor principal: los diferentes tipos de promoción P1, P2, P3, luego $K = 3$

Análisis de la Investigación Cuantitativa

Unidades experimentales: 15 (5 sucursales x 3 tipos de promoción)

Variable dependiente: unidades vendidas

Hipótesis nula (H_0) Los resultados obtenidos son independientes del tipo de promoción, es decir las medias de ventas de las tres promociones serán iguales.

$$H_0 : m_{p1} = m_{p2} = m_{p3}$$

Número total de mediciones: $n=15$

Número de mediciones por cada tratamiento (promoción) $n_j=5$

x_{ij} = unidades físicas vendidas en cada sucursal

m_j = media de unidades vendidas por tratamiento

m = media total

Partiendo del cuadro de resultados, obtenemos los valores de m_j y m , los cuales son:

	S1	S2	S3	S4	S5	S	m_j
P1	65	50	30	40	65	250	50
P2	30	25	15	20	35	125	25
P3	15	10	10	25	50	110	22

por tanto, $m = 32'333$

Una vez obtenidos estos datos, pasamos a realizar los cálculos de la técnica ANOVA

Dispersión total:

$$D T = \sum_{j=1}^k \sum_{i=1}^n (x_{ij} - m)^2$$

Sustituyendo por los correspondientes valores obtenemos:

$$\begin{aligned} DT &= (65 - 32'3)^2 + (50 - 32'3)^2 + (30 - 32'3)^2 + (40 - 32'3)^2 + (65 - 32'3)^2 + \\ &- 32'3)^2 + (25 - 32'3)^2 + (15 - 32'3)^2 + (20 - 32'3)^2 + (35 - 32'3)^2 + (15 - 32'3)^2 + (10 \\ &- 32'3)^2 + (10 - 32'3)^2 + (25 - 32'3)^2 + (50 - 32'3)^2 = 4.693'333 \end{aligned}$$

Dispersión factorial:

$$DF = \sum_{j=1}^k n_j (m_j - m)^2$$

Sustituyendo obtenemos

$$DF = 5(50 - 32'3)^2 + 5(25 - 32'3)^2 + 5(22 - 32'3)^2 = 2.363'333$$

Dispersión residual:

$$DR = DT - DF \quad \text{Luego } DR = 4.693'33 - 2.363'33 = 2.330$$

Cuadrado medio factorial (CMF):

$$CMF = \frac{DF}{gl} = \frac{DF}{k - 1}$$

Sustituyendo obtenemos $CMF = 1.181'6667$

Cuadrado medio residual (CMR):

$$CMR = \frac{DR}{gl} = \frac{DR}{n - k}$$

Sustituyendo obtenemos $CMR = 194'1667$

Test de la F :

$$F = \frac{CMF}{CMR}$$

Sustituyendo obtenemos $F = 6'0858$

Si buscamos el valor de F en tablas para un nivel del 95% y $gl = 2$ y 12 , obtenemos que $F = 3'89$

Como $6'0858 > 3'89$, existe un efecto significativo de los diferentes tratamientos estudiados para un nivel de confianza del 95%.

La conclusión es que los diferentes tipos de promoción afectan significativamente a la demanda.

Es decir rechazamos la hipótesis nula o de independencia

Análisis de la Investigación Cuantitativa

8.3.2 TEST t DE MEDIAS

El test t de medias es una prueba estadística de contraste de hipótesis que nos permite contrastar si existen diferencias significativas entre dos valores medios.

El proceso es

- Se fija la hipótesis nula
- Se fija el nivel de significación α , normalmente 5% (Error tipo I)
- Determinación del valor observado o medido
- Determinación del valor crítico
- Contraste de hipótesis

Este tipo de test es aplicable en el tratamiento de muestras pequeñas. ($n < 30$ unidades).

EJEMPLO: DIFERENCIA DE MEDIAS, DOS POBLACIONES INDEPENDIENTES

¿Existen sueldos medios diferentes según el sexo del empleado? o ¿Existe diferencia de medias en la experiencia previa según el sexo de los empleados?

Estadísticos del grupo

Sexo del empleado		N	Media	Desviación típ.	Error típ. de la media
Sueldo actual en miles de ptas.	Hombre	217	5426,88	2729,35	185,28
	Mujer	191	3361,17	1041,14	75,33
Experiencia previa (en años)	Hombre	216	1,7870	,6231	4,240E-02
	Mujer	191	1,7359	,5943	4,300E-02

Prueba de muestras independientes

		Prueba de Levene para la igualdad de varianzas		Prueba T para la igualdad de medias						
		F	Sig.	t	gl	Sig. (bilateral)	Diferencia de medias	Error típ. de la diferencia	Intervalo de confianza para la diferencia	
									Inferior	Superior
Sueldo actual en miles de ptas.	Se han asumido varianzas iguales	80,26	,000	9,847	406	,000	2065,71	209,78	1653,33	2478,10
	No se han asumido varianzas iguales			10,328	284,482	,000	2065,71	200,01	1672,02	2459,40
Experiencia previa (en años)	Se han asumido varianzas iguales	1,763	,185	,844	405	,399	5,1E-02	6,057E-02	-7,E-02	,1702
	No se han asumido varianzas iguales			,846	402,675	,398	5,1E-02	6,039E-02	-7,E-02	,1698

9. ANÁLISIS DE MUESTRAS PEQUEÑAS

9.1 INTRODUCCIÓN

En muchas ocasiones, en las investigaciones de mercado no se pueden utilizar muestras grandes. Se utilizan, pues, muestras pequeñas, entendiendo por tales aquéllas que tienen menos de treinta elementos ($n < 30$).

Una muestra se considera grande cuando tiene más de 30 elementos para la media (m) y para la proporción (p). Para la determinación del resto de estadísticos, se considera grande la que tiene más de 100 elementos. En este capítulo consideramos como muestra pequeña la que tiene menos de 30 elementos.

Recordemos que para las muestras de más de treinta elementos ($n > 30$), la mayoría de los estadísticos muestrales se ajustan a la distribución normal o de Gauss. En cambio, por regla general, cuando $n < 30$, la distribución muestral no sigue la distribución normal, sino la denominada t de Student.

9.2 DISTRIBUCIÓN “t” DE STUDENT

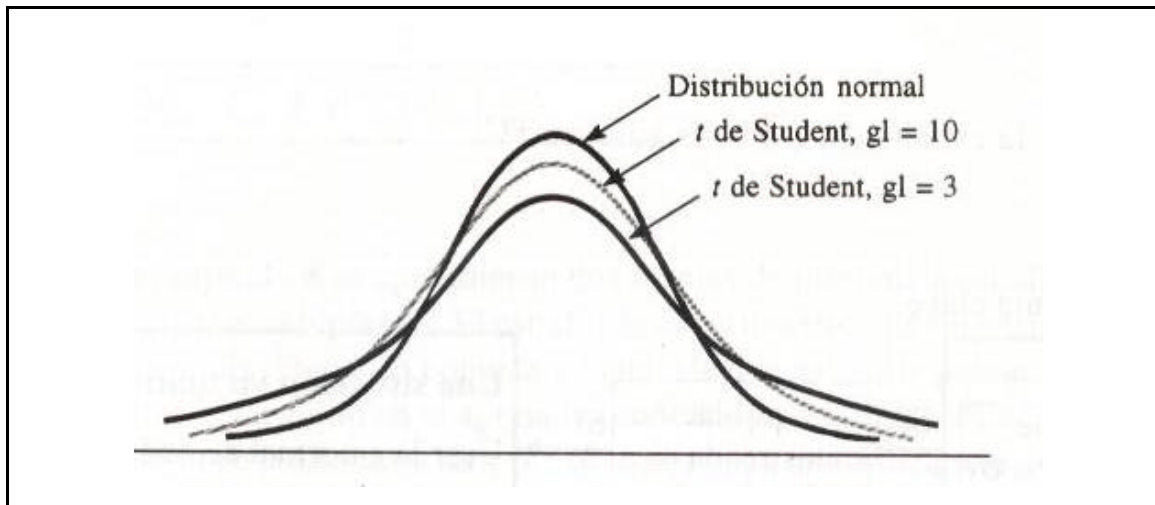
La distribución t de Student es una distribución de mucha dispersión (variabilidad), teniendo más variabilidad cuanto menor es el tamaño (n) de la muestra, o sea, cuantos menos grados de libertad existan. Recordemos que los grados de libertad son $gl = n - 1$. Siendo n el tamaño de la muestra.

Si comparamos la distribución t de Student con la curva normal de Gauss, observaremos que la de Student resulta más baja en el centro y más alta en las colas.

Es decir, en la distribución normal existen más datos alrededor del valor central que en una distribución t de Student, tal como se indica en la figura:⁸

⁸ Robert Jonson y Patricia Kuby Estadística elemental Lo esencial. Editorial Thomson 1998
Investigación Comercial

Fig.: Comparación entre curva normal y t de Student



La distribución t de Student se representa por la siguiente función matemática:

$$Y = \frac{Y_0}{\left(1 + \frac{t^2}{gl}\right)^{\frac{n}{2}}}$$

Donde Y_0 es una ordenada constante que depende del tamaño n de la muestra, gl son los grados de libertad $n - 1$, y t es la razón crítica que viene dada por:

$$t = \frac{\text{estadístico o - parámetro}}{\text{error típico}}$$

Por ejemplo, en el caso de la media vendrá dado por:

$$t = \frac{\frac{m - \mu}{s}}{\frac{s}{\sqrt{n-1}}} = \frac{m - \mu}{s} \times \sqrt{n-1}$$

Donde m es la media muestral, μ es la media de la población o universo, n es el tamaño de la muestra y s es la desviación típica (cuasidesviación típica) de la muestra.

La razón crítica t tiene una distribución diferente según el número de grados de libertad ($n - 1$). Sucediendo que, a partir de 30 grados de libertad, la distribución t es muy similar a la normal.

Para interpretar cada t es preciso consultar la tabla correspondiente de la t de Student, teniendo en cuenta los grados de libertad y el nivel de confianza exigido, para que la razón crítica t obtenida resulte significativa.

9.2.1 FIABILIDAD DE UN ESTADÍSTICO

Para el caso de tratar de determinar la fiabilidad de un estadístico correspondiente a una muestra pequeña, el procedimiento es muy similar al seguido con las muestras que siguen la distribución normal. Vamos a considerar el caso de la media, para un nivel de confianza C. La distribución t de Student también cumple que:

$$m - \mu = t_c \sigma_m$$

siendo:

$$s_m = \frac{s}{\sqrt{gl}}$$

Donde m es la media muestral, s la desviación típica de la muestra, gl los grados de libertad n - 1, μ la media del universo o población y s_m el error típico de la distribución muestral.

El intervalo en el que se encontrará la media del universo para un nivel de confianza C será:

$$m \pm t_c \frac{s}{\sqrt{n-1}}$$

Para un nivel de confianza del 95%, es decir, el 95 % de la medias de las muestras pequeñas, no tendrán una diferencia respecto de la verdadera media de la población mayor que :

$$t_{0'05} * \frac{s}{\sqrt{n-1}}$$

Por consiguiente, dada una media (m) de una muestra, con el 95% de probabilidades de acertar se cumple que:

$$-t_{0'05} < \frac{m - \mu}{s/\sqrt{n-1}} < t_{0'05}$$

De donde se deduce que la media de la población μ se encuentra en el siguiente intervalo

$$m - t_{0'05} \frac{s}{\sqrt{n-1}} < \mu < m + t_{0'05} \frac{s}{\sqrt{n-1}}$$

9.2.2 SIGNIFICACIÓN DE LA MEDIA DE MUESTRAS PEQUEÑAS Y SU FIABILIDAD

Para verificar hipótesis acerca de la verdadera media del universo o simplemente averiguar los intervalos entre los que se encuentra comprendida para un determinado nivel de confianza, el procedimiento que hay que seguir, conociendo la media de una muestra pequeña, es similar al del caso de la distribución normal, salvo que hay que determinar la razón crítica t_c , en vez de la razón crítica Z , consultando este valor t_c en la tabla. Vamos a clarificar lo expuesto con un caso práctico.

CASO PRÁCTICO

Tomamos una muestra de 17 personas. Preguntados acerca de la cuestión, obtenemos los siguientes resultados: para la media y la desviación típica, $m = 50$ y $s = 7$. Queremos saber si la verdadera media para la población puede ser 55, para un nivel de significación del 1 %.

Solución:

En este caso, las hipótesis de trabajo serán: $H_0 : m_1 = m_2$ y la alternativa $H_1 : m \neq 55$

El estadístico de contraste es:

$t = (\text{estadístico} - \text{parámetro}) / \text{error típico}$

El error típico es:

$$s_m = \frac{s}{\sqrt{gl}} = \frac{7}{\sqrt{17-1}} = 1'7$$

Por tanto la razón crítica t será:

$$t = (50 - 55) / 1'7 = - 2'857$$

Consultando en la tabla obtenemos que para 16 grados de libertad y un nivel de confianza de 0'01, el valor de la razón crítica es:

$$t_{0'01} = - 2'583$$

Como el valor obtenido es mayor (en valor absoluto) que el de las tablas

$$2'857 > 2'583 \text{ Rechazamos la } H_0$$

Podemos afirmar con el 99% de probabilidades que no es posible que la media de la población sea 55.

Investigación Comercial

Fiabilidad: la media de la población estará comprendida en: $m \pm t_{0.01} \sigma_m$

Sustituyendo por los valores obtendremos: $50 \pm 2.583 \times 1.7 = 50 \pm 4.52$

Luego la media del universo μ estará comprendida entre 54.52 y 45.48

Es decir $45.48 < \mu < 54.52$ Conclusión no es posible que la media sea 55

9.2.3 SIGNIFICACIÓN DE LA DIFERENCIA DE MEDIAS EN MUESTRAS PEQUEÑAS INDEPENDIENTES

El objetivo de este tipo de estudio es averiguar si la diferencia entre las medias de dos muestras independientes puede ser nula.

El cálculo de la razón crítica viene dado por:

$t = \text{diferencia de medias} / \text{error típico}$

Hay que comparar este valor calculado de la razón crítica con el de las tablas, para el número de grados de libertad correspondiente, que será: $gl = n_1 + n_2 - 2$; donde n_1 y n_2 son el tamaño de las muestras 1 y 2 respectivamente.

El error típico se calcula mediante la siguiente fórmula:

$$s_{m_1 - m_2} = \sqrt{\left[\frac{n_1 s_1^2 + n_2 s_2^2}{gl} \right] * \left[\frac{n_1 + n_2}{n_1 n_2} \right]}$$
$$gl = n_1 + n_2 - 2$$

CASO PRÁCTICO

Realizamos un estudio de mercado sobre dos muestras diferentes, obteniéndose los siguientes resultados:

	Muestra 1	Muestra 2
Tamaño, n	8	14
Media, m	50	45
Error típico, s	5	4

Queremos determinar si existe una diferencia significativa de las medias, para un nivel de significación del 5%.

Solución:

La hipótesis nula será $H_0 : m_1 = m_2$ y la alternativa $H_1 : m_1 \neq m_2$

El número de grados de libertad es:

$$gl = 8 + 14 - 2 = 20$$

Para un nivel α de significación del 5%, el valor de la razón crítica en tablas es

$$t_{0.05} = 1.725.$$

De conformidad con los resultados de la investigación comercial, el valor de la razón crítica t será:

t = diferencia de las medias / error típico

Aplicando la fórmula del error típico obtenemos:

$$s_{1-2} = \sqrt{\left(\frac{8 \cdot 5^2 + 14 \cdot 4^2}{20} \right) \left(\frac{8 + 14}{8 \cdot 14} \right)} = 2.04$$

Luego el valor de t será:

$$t = (50 - 45) / 2.04 = 2.45$$

Como $t \ 2.45 > t_{0.05} \ 1.725$, podemos afirmar que sí que hay una diferencia significativa entre ambas muestras. Es decir, rechazamos la hipótesis nula.

9.2.4 SIGNIFICACIÓN DE LA DIFERENCIA DE MEDIAS DE MUESTRAS PEQUEÑAS RELACIONADAS

Al estar las muestras relacionadas, el error típico de las diferencias entre las medias es menor que en el caso de muestras independientes.

Este tipo de relación suele darse cuando se estudia una misma muestra en tiempos diferentes. Por ejemplo, estudiamos la misma muestra antes y después de una campaña publicitaria (experimento). También pueden ser grupos diferentes, en los que se han emparejado previamente sus componentes, de manera tal que cada uno tiene una réplica casi idéntica en el otro grupo.

Vamos a ver cómo se resuelve esta cuestión mediante la solución de un caso práctico.

CASO PRÁCTICO

Realizamos una investigación de mercado, en una serie de establecimientos comerciales, acerca de la evolución de la venta del producto A, antes y después de realizarse una campaña publicitaria.

Los resultados obtenidos se resumen en la tabla 1.

Análisis de la Investigación Cuantitativa

Se quiere conocer si la media de la segunda medición es significativamente diferente de la media obtenida antes de la campaña, para un nivel de significación del 1%.

Tabla Resultados

Muestra n	Ventas antes de la campaña	Ventas después de la campaña	Diferencia d	$(d - m_d)^2$
1	18	22	4	4
2	25	30	5	9
3	17	15	-2	16
4	30	32	2	0
5	15	17	2	0
6	16	17	1	1
Total	121	133	12	30
Medias	$m_1 = 20'2$	$M_2 = 22'2$	$m_d = 2$	

Solución:

Las hipótesis de este supuesto serán:

$H_0 : m_1 = m_2$ y la alternativa $H_1 : m_1 \neq m_2$

Procederemos realizando los pasos siguientes:

1. Calculamos la desviación típica de las diferencias

$$S_d = \sqrt{\frac{\sum (d - m_d)^2}{n}} = \sqrt{\frac{30}{6}} = 2'24$$

2. Calculamos el error típico de la media de las diferencias

$$s_d = \frac{s_d}{\sqrt{n-1}} = \frac{2'24}{\sqrt{6-1}} = 1$$

3. Determinamos la razón crítica t

$$t = \frac{|m_d|}{s_d} = \frac{2}{1} = 2$$

4. Buscamos en la tabla

Para un nivel de significación del 1% y 5 grados de libertad, obtenemos el siguiente resultado: $t_{0'01} = 3'365$

Como el resultado obtenido es $2 < 3'365$, NO podemos rechazar la hipótesis nula, por consiguiente no podemos afirmar que la diferencia de medias sea significativa.

9.2.5 SIGNIFICACIÓN DE LA DIFERENCIA DE DESVIACIONES TÍPICAS DE MUESTRAS RELACIONADAS

Cuando dos muestras, grandes o pequeñas, están relacionadas y se quiere conocer si sus respectivas desviaciones típicas tienen una diferencia significativa, es necesario determinar la razón crítica t de Student. Se hace mediante la siguiente fórmula:

$$t = \frac{(s_1^2 - s_2^2)\sqrt{n-2}}{2s_1s_2\sqrt{1-r_{12}^2}}$$

Donde n es el tamaño total de las muestras, s_1 y s_2 son las desviaciones típicas de las muestras 1 y 2 respectivamente, y r_{12} es el coeficiente de relación entre ambas muestras. Si la razón crítica t calculada es mayor que la de las tablas, decimos que la diferencia sí es significativa al nivel de confianza utilizado.

CASO PRÁCTICO

Supongamos que trabajamos con una muestra de 51 personas a las que sometemos inicialmente a un experimento, obteniéndose un desviación típica $s_1 = 6$. Un mes después se les somete al mismo experimento y se obtiene una $s_2 = 5$. El coeficiente de relación es de $r_{12} = 0'6$. Nivel de significación del 1%.

Solución: $H_0 : s_1 = s_2$ y la alternativa $H_1 : s_1 \neq s_2$

$$t = \frac{(6^2 - 5^2)\sqrt{51-2}}{2*6*5\sqrt{1-0'6^2}} = 1'604$$

El valor de t para 49 gl y el 1% es, aproximadamente, 2'42.

Como $1'604 < 2'42$, no podemos rechazar H_0 . No hay una diferencia significativa.

9.3 SIGNIFICACIÓN Y FIABILIDAD DEL COEFICIENTE DE CORRELACIÓN R DE PEARSON PARA MUESTRAS PEQUEÑAS

Para determinar la fiabilidad del coeficiente r de Pearson en muestras pequeñas, necesitamos conocer el error típico de la distribución muestral de la r de Pearson.

Las fórmulas que se van a utilizar son:

$$s_r = \sqrt{\frac{1-r^2}{n-2}}$$

$$r_N = r \pm ts_r$$

Si $r > 0.8$ la distribución es asimétrica y se trabaja en base a la Z de Fisher, para trabajar de esta forma con las características de la normal. El error típico en este caso vendrá dado por:

$$s_z = \frac{1}{\sqrt{n-3}}$$

CASO PRÁCTICO

Tenemos 20 pares de datos ($n = 20$), entre los cuales existe una correlación $r = 0.6$. Hay que determinar la fiabilidad de este coeficiente a un nivel de significación del 5%.

Solución:

Mirando en las tablas para un nivel del 5% y $gl = 20 - 2 = 18$ obtenemos $t = 1.734$

$$s_r = \sqrt{\frac{1-r^2}{n-2}} = \sqrt{\frac{1-0.36}{18}} = 0.189$$

$$r_N = r \pm ts_r = 0.6 \pm 1.734 \cdot 0.189 = 0.6 \pm 0.33$$

La variabilidad del coeficiente está entre 0.93 y 0.27; como puede observarse es muy amplia, esto nos indica que es recomendable trabajar con una muestra grande.

9.3.1 SIGNIFICACIÓN DE LA DIFERENCIA ENTRE COEFICIENTES DE CORRELACIÓN OBTENIDOS EN MUESTRAS RELACIONADAS

Cuando entre dos muestras o entre variables de una misma muestra existe una correlación, la razón crítica se calcula de acuerdo con la siguiente fórmula:

$$t = \frac{(r_{12} - r_{13})\sqrt{n-3}\sqrt{1+r_{23}}}{\sqrt{2}\sqrt{1-r_{12}^2-r_{13}^2-r_{23}^2+2r_{12}r_{13}r_{23}}}$$

CASO PRÁCTICO

En un estudio sobre una muestra de 50 personas, encontramos una correlación entre el sexo y la lectura de un determinado periódico, con valor $r = 0.72$, y entre el sexo y la intención de voto a un determinado partido político, $r = 0.78$, si la correlación entre el periódico y el voto al partido político es $r = 0.6$. Queremos conocer si hay diferencia significativa entre $r = 0.72$ y $r = 0.78$, a un nivel de significación del 5%

Solución:

El valor de la razón crítica t será

$$t = \frac{(0.72 - 0.78)\sqrt{50-3}\sqrt{1+0.6}}{\sqrt{2}\sqrt{1-0.72^2-0.78^2-0.6^2+2*0.72*0.78*0.6}} = 1.958$$

En tablas para 47 grados de libertad y 5% obtenemos $t_c = 1.67$

Como $1.958 > 1.67$ deberemos rechazar la hipótesis nula a un nivel del 5%, es decir, tenemos el 95% de probabilidades de que la diferencia entre los dos coeficientes no sea cero. Por consiguiente $r = 0.78$ es mayor significativamente que $r = 0.72$.

Análisis de la Investigación Cuantitativa

TABLA ESTADÍSTICA: DISTRIBUCIÓN t DE STUDENT

Valores de la función de distribución

g.l. = grados de libertad

t_c tal que $p(t \leq t_c) = p$

Probabilidad p										
g.l.	0,995	0,990	0,975	0,950	0,900	0,800	0,750	0,700	0,600	0,550
1	63,657	31,821	12,706	6,314	3,078	1,376	1,000	0,727	0,325	0,158
2	9,925	6,965	4,303	2,920	1,876	1,061	0,816	0,617	0,289	0,142
3	5,841	4,451	3,183	2,353	1,638	0,978	0,765	0,584	0,277	0,137
4	4,604	3,747	2,786	2,132	1,533	0,941	0,741	0,569	0,271	0,134
5	4,032	3,365	2,571	2,015	1,478	0,920	0,727	0,559	0,267	0,132
6	3,707	3,143	2,457	1,943	1,440	0,906	0,718	0,553	0,265	0,131
7	3,499	2,998	2,365	1,895	1,415	0,896	0,711	0,549	0,263	0,130
8	3,355	2,895	2,306	1,860	1,397	0,889	0,706	0,546	0,262	0,130
9	3,250	2,821	2,262	1,833	1,383	0,883	0,703	0,543	0,261	0,129
10	3,169	2,764	2,228	1,812	1,372	0,879	0,700	0,542	0,260	0,129
11	3,106	2,728	2,201	1,796	1,363	0,876	0,697	0,540	0,260	0,129
12	3,055	2,681	2,179	1,782	1,356	0,873	0,695	0,539	0,259	0,128
13	3,012	2,650	2,160	1,771	1,350	0,870	0,694	0,538	0,259	0,128
14	2,987	2,624	2,145	1,761	1,345	0,868	0,692	0,537	0,258	0,128
15	2,947	2,602	2,131	1,753	1,341	0,866	0,691	0,536	0,258	0,128
16	2,921	2,583	2,120	1,746	1,337	0,865	0,690	0,535	0,258	0,128
17	2,898	2,567	2,110	1,740	1,333	0,863	0,689	0,534	0,257	0,128
18	2,888	2,552	2,101	1,734	1,330	0,862	0,688	0,534	0,257	0,127
19	2,861	2,539	2,093	1,729	1,328	0,861	0,688	0,533	0,257	0,127
20	2,845	2,528	2,086	1,725	1,325	0,860	0,687	0,533	0,257	0,127
21	2,831	2,518	2,080	1,721	1,323	0,859	0,686	0,532	0,257	0,127
22	2,819	2,508	2,074	1,717	1,321	0,858	0,686	0,532	0,256	0,127
23	2,807	2,500	2,069	1,714	1,319	0,858	0,685	0,532	0,256	0,127
24	2,797	2,492	2,064	1,711	1,318	0,857	0,685	0,531	0,256	0,127
25	2,787	2,485	2,060	1,708	1,316	0,856	0,684	0,531	0,256	0,127
26	2,779	2,479	2,056	1,706	1,315	0,856	0,684	0,531	0,256	0,127
27	2,771	2,473	2,052	1,703	1,314	0,855	0,684	0,531	0,256	0,127
28	2,763	2,467	2,048	1,701	1,313	0,855	0,683	0,530	0,256	0,127
29	2,756	2,462	2,045	1,699	1,311	0,854	0,683	0,530	0,256	0,127
30	2,750	2,457	2,042	1,697	1,310	0,854	0,683	0,530	0,256	0,127
40	2,704	2,423	2,021	1,684	1,303	0,851	0,681	0,529	0,255	0,126
60	2,660	2,390	2,000	1,671	1,296	0,848	0,679	0,527	0,254	0,126

9.4 DISTRIBUCIÓN CHI CUADRADO (c^2)

Si consideramos una población normal N y tomamos las variables $x_1, x_2, x_3, \dots, x_n$ con media μ y desviación típica σ , es decir $N(\mu; \sigma)$, si tomamos muchísimas muestras de tamaño n y desviación típica s y calculamos para cada una de ellas el valor de:

$$c^2 = \frac{ns^2}{s^2}$$

obtenemos una distribución muestral de χ^2 .

El aspecto más interesante de esta distribución es constatar su dependencia de los grados de libertad y su utilidad para resolver problemas acerca de la desviación típica σ de la población, conociendo la desviación típica s de la muestra pequeña.

Para interpretar un determinado valor de χ^2 acudimos a las tablas correspondientes, teniendo en cuenta los grados de libertad y el nivel de confianza. Los grados de libertad son $n - 1$. Si s es menor que σ habrá que mirar a la izquierda, y si $s > \sigma$, habrá que mirar a la derecha.

CASO PRÁCTICO 1

De un universo normal con desviación típica 14 se toma una muestra de 20 elementos y se obtiene una desviación típica de 17. ¿Está bien tomada la muestra?

Solución:

Aplicando la fórmula

$$c^2 = \frac{ns^2}{s^2} = \frac{20 \cdot 17^2}{14^2} = 29'49$$

Tomamos en tablas para 19 grados de libertad y significación del 5% y obtenemos el valor 30'144. Como este valor es mayor que el calculado, podemos pensar con un 95% de probabilidades que la muestra está bien tomada.

CASO PRÁCTICO 2

Una muestra de 25 elementos tomada de una población normal, nos da una desviación típica de $s = 15$. ¿En qué intervalo encontraremos la desviación típica de la población? Intervalo de confianza 96%, es decir 2% a cada lado de la distribución χ^2

Solución:

Para el intervalo objeto de estudio se cumplirá que:

$$c_{0.98}^2 \leq \frac{ns^2}{s^2} \leq c_{0.02}^2$$
$$\frac{ns^2}{c_{0.98}^2} \geq s^2 \geq \frac{ns^2}{c_{0.02}^2}$$

Sustituyendo:

$$\sqrt{\frac{25 * 15^2}{11.992}} \geq s \geq \sqrt{\frac{25 * 15^2}{40.270}}$$

Luego podemos afirmar que la desviación típica de la población, con un error del 4%, se encuentra entre 21.657 y 11.818.

Análisis de la Investigación Cuantitativa

TABLA ESTADÍSTICA: DISTRIBUCIÓN DE χ^2

Valores de la función de distribución

g.l. = grados de libertad

χ^2_c tal que $p(\chi^2 \leq \chi^2_c) = p$

Probabilidad p

g.l.	0,995	0,990	0,975	0,950	0,900	0,500	0,100	0,050	0,025	0,010	0,005
1	7,88	6,63	5,02	3,84	2,71	0,45	0,01	0,00	0,00	0,00	0,00
2	10,60	9,21	7,38	5,99	4,61	1,39	0,21	0,10	0,05	0,02	0,01
3	12,84	11,34	9,35	7,81	6,25	2,37	0,58	0,35	0,22	0,12	0,07
4	14,86	13,28	11,14	9,49	7,78	3,36	1,06	0,71	0,48	0,30	0,21
5	16,75	15,09	12,83	11,17	9,24	4,25	1,61	1,15	0,83	0,55	0,41
6	18,55	16,81	14,45	12,69	10,64	5,35	2,20	1,64	1,24	0,87	0,68
7	20,28	18,48	16,01	14,07	12,02	6,35	2,83	2,17	1,69	1,24	0,99
8	21,96	20,09	17,53	15,51	13,36	7,34	3,49	2,73	2,18	1,65	1,34
9	23,59	21,67	19,02	16,92	14,68	8,34	4,17	3,33	2,70	2,09	1,73
10	25,19	23,21	20,48	18,31	15,99	9,34	4,87	3,94	3,25	2,56	2,16
11	26,76	24,73	21,92	19,68	17,28	10,34	5,58	4,57	3,82	3,05	2,60
12	28,30	26,22	23,34	21,03	18,55	11,34	6,30	5,23	4,40	3,57	3,07
13	29,82	27,69	24,74	22,36	19,81	12,34	7,04	5,89	5,01	4,11	3,57
14	31,32	29,14	26,12	23,68	21,06	13,34	7,79	6,57	5,63	4,66	4,07
15	32,80	30,58	27,49	25,00	22,31	14,34	8,55	7,26	6,26	5,23	4,60
16	34,27	32,00	28,85	26,30	23,54	15,34	9,31	7,96	6,91	5,81	5,14
17	35,72	33,41	30,19	27,59	24,77	16,34	10,09	8,67	7,56	6,41	5,70
18	37,16	34,81	31,53	28,87	25,99	17,34	10,86	9,39	8,23	7,01	6,26
19	38,58	36,29	32,85	30,14	27,20	18,34	11,65	10,12	8,91	7,63	6,84
20	40,00	37,67	34,27	31,41	28,41	19,34	12,44	10,85	9,59	8,26	7,43
21	41,40	38,93	35,48	32,67	29,62	20,34	13,24	11,59	10,28	8,90	8,03
22	42,80	40,29	36,78	33,92	30,81	21,34	14,04	12,34	10,98	9,54	8,64
23	44,18	41,64	38,08	35,17	32,01	22,34	14,85	13,09	11,69	10,20	9,26
24	45,56	42,98	39,36	36,42	33,20	23,34	15,66	13,85	12,40	10,86	9,89
25	46,93	44,31	40,65	37,65	34,38	24,34	16,47	14,61	13,12	11,52	10,52
26	48,29	45,64	41,92	38,89	35,56	25,34	17,29	15,38	13,84	12,20	11,16
27	49,64	46,96	43,29	40,11	36,74	26,34	18,11	16,15	14,57	12,83	11,81
28	50,99	48,28	44,46	41,34	37,92	27,34	18,94	16,93	15,31	13,56	12,46
29	52,34	49,59	45,72	42,56	39,09	28,34	19,77	17,71	16,05	14,26	13,12
30	53,67	50,89	46,98	43,77	40,26	29,34	20,60	18,49	16,89	14,96	13,78
40	66,77	63,69	59,34	55,76	51,81	39,34	29,05	26,51	24,43	22,16	20,71
60	91,95	88,38	83,30	79,08	74,40	59,34	46,56	43,19	40,48	37,43	35,58

9.5 DISTRIBUCIÓN F DE FISHER

Cuando estudiamos la significación de la diferencia de las desviaciones típicas de dos muestras pequeñas independientes obtenidas de una misma población, podemos extraer consecuencias averiguando si las diferencias entre las desviaciones típicas pueden ser 0 por azar o no, en cuyo caso concluiremos que una muestra posee características propias que la hacen ser diferente de la otra.

La distribución F de Fisher es la razón entre dos estimaciones de la varianza de la población y sirve para poner a prueba el hecho de que dos estimaciones lo sean realmente de una varianza igual.

El proceso consiste en saber si una razón F, calculada entre dos muestras determinadas, es compatible con la hipótesis de que ambas muestras procedan de una única población y sólo difieran entre sí por azar.

Una vez calculada la F es necesario acudir a las correspondientes tablas, teniendo en cuenta los grados de libertad correspondientes a cada muestra ($n - 1$) y el nivel de significación, normalmente el 5%.

Si la F calculada resulta mayor que la F de tablas a un nivel de significación del 5%, podemos afirmar que F no puede darse ni en el 5% de los casos; es decir, las dos muestras son significativamente diferentes, no pudiendo resultar iguales por azar ni en el 5% de los casos.

Para el cálculo se utiliza la siguiente fórmula:

$$F = \frac{s_1^2}{s_2^2} = \frac{\frac{\sum (x_i - m_1)^2}{n_1 - 1}}{\frac{\sum (x_i - m_2)^2}{n_2 - 1}}$$

En el numerador se pone siempre la varianza mayor.

CASO PRÁCTICO

Los resultados correspondientes a dos muestras pequeñas son los que se recogen en la tabla siguiente. ¿Pertenecen ambas muestras al mismo universo? Nivel de significación $\alpha = 5\%$.

n_1	n_2
20	19
17	20
16	18
14	30
20	25
18	
27	
$\Sigma = 132$	$\Sigma = 112$
$m_1 = 18'86$	$m_2 = 22'4$
$(s_1)^2 = 17'5$	$(s_2)^2 = 26'9$

Los valores de m y s se han aproximado a dos y un decimal.

Si ambas muestras pertenecen a un mismo universo, las estimaciones respectivas de la varianza no podrán ser significativamente diferentes.

El valor de F de tablas, para un nivel de confianza del 5% y siendo los grados de libertad 4 y 6, obtenemos que: $F_t = 4'53$

El valor de F en este estudio es de

$$F = \frac{s_2^2}{s_1^2} = \frac{26'9}{17'5} = 1'537 = 1'54$$

Como 1'54 es menor que 4'53, podemos decir que ambas muestras corresponden a una misma población.

Análisis de la Investigación Cuantitativa

TABLA ESTADÍSTICA: DISTRIBUCIÓN DE LA F
NIVEL DE CONFIANZA 95%

n	m				
	1	2	3	4	5
1	161´4	199´5	215´7	224´6	230´2
2	18´51	19	19´16	19´25	19´30
3	10´13	9´55	9´28	9´12	9´01
4	7´71	6´94	6´59	6´39	6´26
5	6´61	5´79	5´41	5´19	5´05
6	5´99	5´14	4´76	4,53	4´39
7	5´59	4´74	4´35	4´12	3´97
8	5´32	4´46	4´07	3´84	3´69
9	5´12	4´26	3´86	3´63	3´48
10	4´96	4´10	3´71	3´48	3´33
11	4´84	3´98	3´59	3´36	3´20
12	4´75	3´89	3´49	3´26	3´11
13	4´67	3´81	3´41	3´18	3´03
14	4´6	3´74	3´34	3´11	2´96
15	4´54	3´68	3´29	3´06	2´90

Siendo **m** los grados de libertad del numerador y **n** los grados de libertad del denominador.

10. TEST PARAMÉTRICOS

10.1 INTRODUCCIÓN

Los test o pruebas paramétricas son técnicas que se fundamentan en varios supuestos que casi nunca son confirmados, debido, fundamentalmente, a que se desconocen las características de la población -o universo- de la que se obtienen los datos. Estos supuestos o condiciones son:

- Las variables consideradas de la población siguen una distribución estadística determinada. Suele ser la norma.
- Las variables objeto de estudio deben estar medidas al menos en escala de intervalo
- El tamaño muestral ha de ser suficientemente grande , es decir más de 30 unidades
- En el caso de dos poblaciones, estas han de tener varianzas similares
- Las muestras normalmente son independientes
- En caso de muestras pareadas, se utiliza el test “t” para pareados

10.2 TIPOS DE TEST PARAMÉTRICOS

Se distingue entre las siguientes modalidades de test paramétricos:

10.2.1 CONTRASTES PARA UNA MUESTRA

Test Z. Está basado en la distribución normal y se utiliza para contrastar estadísticos de una distribución con respecto a valores de parámetros del universo. El paquete estadístico calcula el valor de Z y luego calcula la probabilidad asociada al mismo. Se compara la probabilidad con el nivel de significación seleccionado, no se rechaza la H_0 en el caso de que la probabilidad sea mayor que dicho nivel. Se acepta la hipótesis alternativa cuando la probabilidad es menor que el nivel de confianza.

Test t. Se basa en la t de Student. Parte del conocimiento de la media de la muestra y de la desviación típica de la media y de la media del universo. La interpretación de los resultados es idéntica al test Z.

10.2.2 CONTRASTES PARA DOS MUESTRAS INDEPENDIENTES

Test t para dos grupos independientes. Se aplica a variables medidas en escalas de intervalo o de razón. Se utiliza para contrastar si dos muestras independientes proceden de poblaciones con la misma media.

Test F (F de Barlett Box). Se utiliza para contrastar si las muestras proceden de poblaciones con idénticas varianzas. Realiza el cociente entre las medias cuadráticas de las dos muestras, teniendo en cuenta el número de grados de libertad del numerador y del denominador, asignando así la probabilidad.

10.2.3 CONTRASTES PARA DOS MUESTRAS RELACIONADAS

Test t para pareado. Se utiliza para comprobar si dos muestras provienen de poblaciones con igual media.

10.2.4 PRUEBAS MÁS UTILIZADAS

Las pruebas más utilizadas son: test de la media, test de diferencias de medias en muestras independientes, test de proporciones, test de diferencia de proporciones y test de muestras relacionadas. A continuación explicaremos cada uno de ellos.

10.3 TEST DE LA MEDIA

Se utiliza para comparar la media obtenida de una muestra con la que se presupone de la población. También se utilizar para determinar si dos muestras han sido extraídas de la misma población. Se utilizan las pruebas Z y t.

10.3.1 PRUEBA Z

Esta prueba es adecuada en los siguientes casos:

La distribución del universo es normal, conocemos la varianza, la muestra puede ser pequeña o grande.

La distribución del universo es normal, desconocemos la varianza, la muestra es grande (+ de 30 elementos).

La distribución no es normal, pero el tamaño de muestra es lo suficientemente grande. Aplicando el teorema central del límite la distribución se aproxima a la normal.

El estadístico a utilizar tiene la siguiente fórmula:

$$Z = \frac{m - m}{s}$$

Donde **m** es la media de la muestra, **m** es la media de la población y **s** es la desviación típica del universo. Cuando no conocemos la varianza de la población, la desviación típica se estima mediante la siguiente fórmula:

$$s = \frac{S}{\sqrt{n}}$$
$$S = \sqrt{\frac{\sum_{i=1}^n (x_i - m)^2}{n - 1}}$$

Donde **n** es el tamaño de la muestra **s** es la desviación típica de la media muestral **m** es la media de la muestra, **m** es la media de la población y **s** es la desviación típica del universo (**n - 1** son los grados de libertad).

El valor teórico de Z es para el 5% de significación 1'96 y para el 1% es 2'58

10.3.2 PRUEBA t

Esta prueba es adecuada en los siguientes casos:

La distribución del universo o población es normal

La varianza es desconocida

El tamaño de muestra es pequeño. Si es grande, se aplica la prueba Z

El estadístico a utilizar tiene la siguiente fórmula:

$$t = \frac{m - m}{S_m}$$

con n - 1 grados de libertad

El valor teórico tabulado del estadístico t se recoge en la tabla siguiente:

Análisis de la Investigación Cuantitativa

TABLA ESTADÍSTICA: DISTRIBUCIÓN t DE STUDENT

Valores de la función de distribución

g.l. = grados de libertad

t_c tal que $p(t \leq t_c) = p$

g.l.	Probabilidad p									
	0,995	0,990	0,975	0,950	0,900	0,800	0,750	0,700	0,600	0,550
1	63,657	31,821	12,706	6,314	3,078	1,376	1,000	0,727	0,325	0,158
2	9,925	6,965	4,303	2,920	1,876	1,061	0,816	0,617	0,289	0,142
3	5,841	4,451	3,183	2,353	1,638	0,978	0,765	0,584	0,277	0,137
4	4,604	3,747	2,786	2,132	1,533	0,941	0,741	0,569	0,271	0,134
5	4,032	3,365	2,571	2,015	1,478	0,920	0,727	0,559	0,267	0,132
6	3,707	3,143	2,457	1,943	1,440	0,906	0,718	0,553	0,265	0,131
7	3,499	2,998	2,365	1,895	1,415	0,896	0,711	0,549	0,263	0,130
8	3,355	2,895	2,306	1,860	1,397	0,889	0,706	0,546	0,262	0,130
9	3,250	2,821	2,262	1,833	1,383	0,883	0,703	0,543	0,261	0,129
10	3,169	2,764	2,228	1,812	1,372	0,879	0,700	0,542	0,260	0,129
11	3,106	2,728	2,201	1,796	1,363	0,876	0,697	0,540	0,260	0,129
12	3,055	2,681	2,179	1,782	1,356	0,873	0,695	0,539	0,259	0,128
13	3,012	2,650	2,160	1,771	1,350	0,870	0,694	0,538	0,259	0,128
14	2,987	2,624	2,145	1,761	1,345	0,868	0,692	0,537	0,258	0,128
15	2,947	2,602	2,131	1,753	1,341	0,866	0,691	0,536	0,258	0,128
16	2,921	2,583	2,120	1,746	1,337	0,865	0,690	0,535	0,258	0,128
17	2,898	2,567	2,110	1,740	1,333	0,863	0,689	0,534	0,257	0,128
18	2,888	2,552	2,101	1,734	1,330	0,862	0,688	0,534	0,257	0,127
19	2,861	2,539	2,093	1,729	1,328	0,861	0,688	0,533	0,257	0,127
20	2,845	2,528	2,086	1,725	1,325	0,860	0,687	0,533	0,257	0,127
21	2,831	2,518	2,080	1,721	1,323	0,859	0,686	0,532	0,257	0,127
22	2,819	2,508	2,074	1,717	1,321	0,858	0,686	0,532	0,256	0,127
23	2,807	2,500	2,069	1,714	1,319	0,858	0,685	0,532	0,256	0,127
24	2,797	2,492	2,064	1,711	1,318	0,857	0,685	0,531	0,256	0,127
25	2,787	2,485	2,060	1,708	1,316	0,856	0,684	0,531	0,256	0,127
26	2,779	2,479	2,056	1,706	1,315	0,856	0,684	0,531	0,256	0,127
27	2,771	2,473	2,052	1,703	1,314	0,855	0,684	0,531	0,256	0,127
28	2,763	2,467	2,048	1,701	1,313	0,855	0,683	0,530	0,256	0,127
29	2,756	2,462	2,045	1,699	1,311	0,854	0,683	0,530	0,256	0,127
30	2,750	2,457	2,042	1,697	1,310	0,854	0,683	0,530	0,256	0,127
40	2,704	2,423	2,021	1,684	1,303	0,851	0,681	0,529	0,255	0,126
60	2,660	2,390	2,000	1,671	1,296	0,848	0,679	0,527	0,254	0,126

EJEMPLOS:

Caso 1

Un panadero quiere lanzar un nuevo producto. Para que el proyecto sea viable, requiere de unas ventas mínimas de 80 Kg. semanales por establecimiento. Realiza un estudio de mercado seleccionando 10 establecimientos representativos de su clientela. Los resultados del estudio son los relacionados en la siguiente tabla.

Cliente	Consumo en Kg	Cliente	Consumo en Kg.
1	130	6	70
2	118	7	120
3	60	8	97
4	75	9	150
5	100	10	90

Suponemos que la venta se distribuye de acuerdo con la normal. Trabajamos con un nivel de significación del 5% ¿Cuál sería tu recomendación?

Solución⁹

Desconocemos la varianza del universo y se trata de una muestra pequeña, por lo que el test apropiado sería el t.

Hipótesis

La hipótesis nula será $H_0 : \mu \leq 80$ y la hipótesis alternativa será $H_1 : \mu > 80$

El estadístico a calcular es:

$$t = \frac{m - m}{S_m}$$

con $n - 1$ grados de libertad

Sustituyendo obtenemos

⁹ En los cálculos utilizamos dos decimales. pasando de cinco aumentamos una unidad.

Análisis de la Investigación Cuantitativa

El tamaño de la muestra es $n = 10$, luego los grados de libertad serán $gl = 10 - 1 = 9$

La media muestral será

$$m = \frac{\sum_{i=1}^n x_i}{n} = \frac{130 + 118 + 60 + 75 + 100 + 70 + 120 + 97 + 150 + 90}{10} = 101$$

La media del universo será $\mu = 80$

La desviación típica es

$$s = \frac{S}{\sqrt{n}}$$
$$S = \sqrt{\frac{\sum_{i=1}^n (x_i - m)^2}{n - 1}}$$

Sustituyendo obtenemos que $S = 28'57$ y $\sigma = 9'04$

El valor del estadístico t será

$$t = \frac{m - \mu}{S_m} = \frac{101 - 80}{9'04} = 2'323$$

Para un nivel de significación del 5% y 9 grados de libertad, el valor de tablas de t_α es de 1'833. Como el valor calculado 2'323 es mayor que el teórico de tablas, se rechaza la hipótesis nula; por lo tanto, el proyecto es viable.

Caso 2

Supongamos que el panadero realiza el estudio con una muestra de 1100 establecimientos, obteniendo una media muestral $m = 97$ y una $S_m = 9'2$.

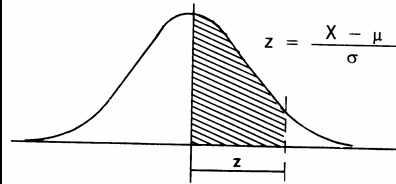
En este caso, se trata de una muestra grande en la que desconocemos la varianza del universo. Se debe aplicar, por tanto, el test Z .

$$Z = \frac{m - \mu}{s} = \frac{97 - 80}{9'2} = 1'8478 = 1'85$$

Como el valor obtenido es menor que el teórico $Z = 1'96$, no rechazaríamos la H_0 . En principio y salvo otros criterios, el proyecto no es viable.

Análisis de la Investigación Cuantitativa

DISTRIBUCIÓN NORMAL TIPIFICADA



z	0,00	0,01	0,02	0,03	0,04	0,05	0,06	0,07	0,08	0,09
0,0	0,0000	0,0040	0,0080	0,0120	0,0160	0,0199	0,0239	0,0279	0,0319	0,0359
0,1	0,0398	0,0438	0,0478	0,0517	0,0557	0,0596	0,0636	0,0675	0,0714	0,0753
0,2	0,0793	0,0832	0,0871	0,0910	0,0948	0,0987	0,1026	0,1064	0,1103	0,1141
0,3	0,1179	0,1217	0,1255	0,1293	0,1331	0,1368	0,1406	0,1443	0,1480	0,1517
0,4	0,1554	0,1591	0,1628	0,1664	0,1700	0,1736	0,1772	0,1808	0,1844	0,1879
0,5	0,1915	0,1950	0,1985	0,2019	0,2054	0,2088	0,2123	0,2157	0,2190	0,2224
0,6	0,2257	0,2291	0,2324	0,2356	0,2389	0,2422	0,2454	0,2486	0,2518	0,2549
0,7	0,2580	0,2611	0,2642	0,2673	0,2704	0,2734	0,2764	0,2793	0,2823	0,2852
0,8	0,2881	0,2910	0,2939	0,2967	0,2996	0,3023	0,3051	0,3078	0,3106	0,3133
0,9	0,3159	0,3186	0,3212	0,3238	0,3264	0,3289	0,3315	0,3340	0,3365	0,3389
1,0	0,3413	0,3437	0,3461	0,3485	0,3508	0,3531	0,3554	0,3577	0,3599	0,3621
1,1	0,3643	0,3665	0,3686	0,3708	0,3749	0,3770	0,3790	0,3810	0,3830	0,3849
1,2	0,3869	0,3888	0,3907	0,3925	0,3944	0,3962	0,3980	0,3997	0,4015	0,4032
1,3	0,4049	0,4066	0,4082	0,4099	0,4115	0,4131	0,4147	0,4162	0,4177	0,4192
1,4	0,4207	0,4222	0,4236	0,4251	0,4265	0,4279	0,4292	0,4306	0,4319	0,4332
1,5	0,4345	0,4357	0,4370	0,4382	0,4394	0,4406	0,4418	0,4429	0,4441	0,4452
1,6	0,4463	0,4474	0,4484	0,4495	0,4505	0,4515	0,4525	0,4535	0,4545	0,4554
1,7	0,4564	0,4573	0,4582	0,4591	0,4599	0,4608	0,4616	0,4625	0,4633	0,4641
1,8	0,4648	0,4656	0,4664	0,4671	0,4678	0,4686	0,4693	0,4699	0,4706	0,4713
1,9	0,4719	0,4726	0,4732	0,4738	0,4744	0,4750	0,4756	0,4761	0,4767	0,4772
2,0	0,4778	0,4783	0,4788	0,4793	0,4798	0,4803	0,4808	0,4812	0,4817	0,4821
2,1	0,4826	0,4830	0,4834	0,4838	0,4842	0,4846	0,4850	0,4854	0,4857	0,4861
2,2	0,4864	0,4868	0,4871	0,4875	0,4878	0,4881	0,4884	0,4887	0,4890	0,4893
2,3	0,4896	0,4898	0,4901	0,4904	0,4906	0,4909	0,4911	0,4913	0,4916	0,4918
2,4	0,4920	0,4922	0,4925	0,4927	0,4929	0,4931	0,4932	0,4934	0,4936	0,4938
2,5	0,4940	0,4941	0,4943	0,4945	0,4946	0,4948	0,4949	0,4951	0,4952	0,4953
2,6	0,4955	0,4956	0,4957	0,4959	0,4960	0,4961	0,4962	0,4963	0,4964	0,4965
2,7	0,4966	0,4967	0,4968	0,4969	0,4970	0,4971	0,4972	0,4973	0,4974	0,4974
2,8	0,4975	0,4976	0,4977	0,4977	0,4978	0,4979	0,4979	0,4980	0,4981	0,4981
2,9	0,4981	0,4982	0,4982	0,4983	0,4984	0,4984	0,4985	0,4985	0,4986	0,4986
3,0	0,4987	0,4987	0,4987	0,4988	0,4988	0,4989	0,4989	0,4989	0,4990	0,4990
3,1	0,4990	0,4991	0,4991	0,4991	0,4992	0,4992	0,4992	0,4992	0,4993	0,4993
3,2	0,4993	0,4993	0,4994	0,4994	0,4994	0,4994	0,4994	0,4995	0,4995	0,4995
3,3	0,4995	0,4995	0,4995	0,4996	0,4996	0,4996	0,4996	0,4996	0,4996	0,4997
3,4	0,4997	0,4997	0,4997	0,4997	0,4997	0,4997	0,4997	0,4997	0,4997	0,4998
3,5	0,4998	0,4998	0,4998	0,4998	0,4998	0,4998	0,4998	0,4998	0,4998	0,4998
3,6	0,4998	0,4998	0,4998	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999
3,7	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999
3,8	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999
3,9	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999	0,4999
4,0	0,5000	0,5000	0,5000	0,5000	0,5000	0,5000	0,5000	0,5000	0,5000	0,5000

10.3.3 PRUEBA DE DIFERENCIA DE MEDIAS INDEPENDIENTES.

A su vez se divide en función de si la varianza del universo es conocida o desconocida

10.3.3.1 CON VARIANZA CONOCIDA

Si se conoce la varianza por algún estudio previo y la distribución es normal, se aplicará la prueba Z si la muestra es grande y el test t si la muestra es pequeña.

El estadístico correspondiente es

$$Z = \frac{m_a - m_b}{s_{m_a - m_b}}$$
$$\text{Siendo } s_{m_a - m_b} = \sqrt{\frac{s_a^2}{n_a} + \frac{s_b^2}{n_b}}$$

Donde m_a y m_b representa las medias de las muestras a y b, n_a y n_b representan el tamaño para las muestras a y b. $s_{m_a - m_b}$ es el error estándar y s_a y s_b las correspondientes desviaciones típicas de la población a y b.

EJEMPLO:

Realizado un estudio sobre 1.000 hogares, se comprueba que la compra media mensual de detergente líquido ha sido de 7.000 cajas, con una desviación típica de 49 cajas.

Se realiza una campaña de comunicación. Se hace un nuevo estudio de mercado, con una muestra de 2.000 hogares, obteniéndose los siguientes resultados: venta media mensual 7.500 cajas con una desviación típica de 15 cajas.

Podemos afirmar que la campaña ha sido eficaz a un nivel del 1%.

Solución:

$$m_1 = 7.500 \quad s_1 = 15 \quad n_1 = 2.000 \quad m_2 = 7.000 \quad s_2 = 49 \quad n_2 = 1.000$$

$$s_{m_1 - m_2} = \sqrt{\frac{15^2}{2000} + \frac{49^2}{1000}} = 1'59$$

$$\frac{\text{Diferencia Observada}}{\text{Error Estándar}} = \frac{7500 - 7000}{1'59} = 314'5$$

Como $314'5 > 2'58$, podemos decir que la campaña ha sido eficaz.

10.3.3.2 CON VARIANZAS DESCONOCIDAS

En este caso, se estima la desviación típica poblacional a partir de las muestrales.

Cuando se trata de una distribución normal pero el tamaño de muestra es pequeño, el estadístico adecuado es el test t con $(n_a + n_b - 1)$ grados de libertad. En el caso de que la muestra sea grande, se utiliza el test Z.

Si la distribución no es normal y la muestra es pequeña se deben aplicar test no paramétricos. Si la muestra es grande por el teorema central del límite se admite la normalidad.

El error estándar en este caso es:

$$S_{m_a - m_b} = \sqrt{s_{m_a}^2 + s_{m_b}^2} = \sqrt{\frac{s_a^2}{n_a} + \frac{s_b^2}{n_b}}$$

Si las varianzas, aún siendo desconocidas, podemos considerarlas iguales, utilizaremos la conjunta. La correspondiente desviación típica será

$$S = \sqrt{\frac{\sum_{i=1}^{n_a} (x_{ia} - m_a)^2 + \sum_{i=1}^{n_b} (x_{ib} - m_b)^2}{n_a + n_b - 2}}$$

Y el error estándar por tanto será

$$S_{m_a - m_b} = \sqrt{\frac{s_a^2}{n_a} + \frac{s_b^2}{n_b}} = \sqrt{S^2 \left(\frac{1}{n_a} + \frac{1}{n_b} \right)}$$

Para decidir si las varianzas son iguales o diferentes se aplica el test de la F con las siguientes hipótesis:

H_0 , las varianzas son iguales y H_1 , las varianzas son diferentes.

El estadístico F se calcula con la siguiente fórmula:

$$F = \frac{s_a^2}{s_b^2}$$

El teórico se obtiene para los siguientes grados de libertad $gl = (n_a - 1) (n_b - 1)$ mirando en las correspondientes tablas. Para calcular el test de las diferencias de medias, si se rechaza la H_0 se utilizan las varianzas separadas; en caso contrario, se utiliza la varianza conjunta. Ver tabla adjunta Distribución F para un nivel de significación del 5%

Análisis de la Investigación Cuantitativa

TABLA ESTADÍSTICA: DISTRIBUCIÓN DE LA F
NIVEL DE CONFIANZA 95%

n	m				
	1	2	3	4	5
1	161´4	199´5	215´7	224´6	230´2
2	18´51	19	19´16	19´25	19´30
3	10´13	9´55	9´28	9´12	9´01
4	7´71	6´94	6´59	6´39	6´26
5	6´61	5´79	5´41	5´19	5´05
6	5´99	5´14	4´76	4,53	4´39
7	5´59	4´74	4´35	4´12	3´97
8	5´32	4´46	4´07	3´84	3´69
9	5´12	4´26	3´86	3´63	3´48
10	4´96	4´10	3´71	3´48	3´33
11	4´84	3´98	3´59	3´36	3´20
12	4´75	3´89	3´49	3´26	3´11
13	4´67	3´81	3´41	3´18	3´03
14	4´6	3´74	3´34	3´11	2´96
15	4´54	3´68	3´29	3´06	2´90

Siendo **m** los grados de libertad del numerador y **n** los grados de libertad del denominador.

10.4 TEST DE PROPORCIONES

Es parecido al anterior de la media, en este caso trabajamos con proporciones.

Los estadísticos a utilizar son:

Para muestras grandes

$$Z = \frac{p - p}{\sqrt{pq/n}}$$

Para muestras pequeñas

$$t = \frac{p - p}{\sqrt{pq/n}}$$

Donde

n es el tamaño de la muestra, **p** es la proporción observada en la muestra, **q = 1 - p** y **p** es la proporción en la población o la de comparación.

10.5 TEST DE SIGNIFICACIÓN PARA POBLACIONES INFINITAS Y QUE SIGUEN LA DISTRIBUCIÓN NORMAL

En Investigación Comercial, los test de significación para muestras grandes de poblaciones infinitas y que siguen la distribución normal más utilizados son:

- test de significación para diferencias de proporciones independientes
- test de significación para diferencias de proporciones no independientes
- test de significación para diferencias de medias de muestras independientes

El estadístico a calcular es el test Z, que viene dado en todos los casos por la fórmula siguiente:

$$\frac{\text{Diferencia Observada}}{\text{Error Estándar}}$$

Para que la diferencia tenga significación al nivel del 5%, el cociente obtenido deberá ser mayor que 1'96. Para que la diferencia sea significativa a un nivel del 1%, el cociente deberá ser mayor que 2'58.

10.5.1 TEST DE SIGNIFICACIÓN PARA DIFERENCIAS DE PROPORCIONES INDEPENDIENTES

Este tipo de prueba se realiza para establecer si la diferencia existente entre dos porcentajes independientes de muestras diferentes, es realmente significativa. Para ello se compara la diferencia existente entre las proporciones y el error estándar de la diferencia de esos porcentajes.

El estadístico que hay que utilizar en este caso concreto es:

$$\frac{\text{Diferencia Observada}}{\text{Error Estándar}} = \frac{p_1 - p_2}{s_{p_1} - s_{p_2}}$$

Donde:

$$s_{p_1} - s_{p_2} = \sqrt{pq \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}$$
$$s_{p_1} - s_{p_2} = \sqrt{\left(\frac{p_1 q_1}{n_1} + \frac{p_2 q_2}{n_2} \right)}$$

Siendo:

$$p = \frac{n_1 p_1 + n_2 p_2}{n_1 + n_2}$$

$$\text{y } q = 100 - p$$

Donde n_1 y n_2 son los respectivos tamaños de las muestras, y p_1 y p_2 son los porcentajes observados.

Las hipótesis serán:

H_0 . No existe diferencia entre las proporciones, $\pi_1 - \pi_2 = 0$

H_1 . Existe desigualdad entre las proporciones, $\pi_1 - \pi_2 \neq 0$

EJEMPLO:

Realizado un estudio sobre una muestra de 1.000 hogares, se observó que el 75% de los mismos carecían de horno microondas. La asociación de fabricantes realizó una campaña de comunicación, con el objetivo de incentivar las ventas. Transcurridos seis meses se realizó una nueva investigación; se tomó una muestra de 800 hogares, obteniéndose como resultado que el 70% de los hogares carecían del citado horno. ¿Es significativa la diferencia observada a un nivel del 5%?

Investigación Comercial

Solución:

$$p_1 = 75 \quad n_1 = 1.000 \quad p_2 = 70 \quad n_2 = 800$$

$$p = \frac{1000 * 75 + 800 * 70}{1000 + 800} = 72'8$$

$$q = 100 - 72'8 = 27'2$$

$$S_{p_1} - S_{p_2} = \sqrt{72'8 * 27'2 \left(\frac{1}{1000} + \frac{1}{800} \right)} = 2'11$$

$$\frac{\text{Diferencia Observada}}{\text{Error Estándar}} = \frac{75 - 70}{2'11} = 2'4$$

Como $2'4 \geq 1'96$, la diferencia encontrada sí que es significativa a un nivel del 5%

10.5.2 TEST DE SIGNIFICACIÓN PARA DIFERENCIAS DE PROPORCIONES NO INDEPENDIENTES

Esta prueba se utiliza para comprobar si la diferencia existente entre dos proporciones relacionadas entre sí es realmente significativa. Los porcentajes pueden estar relacionados de dos formas diferentes:

- Porcentajes excluyentes. Es cuando el elemento pertenece sólo a un determinado porcentaje
- Porcentajes solapados. El elemento pertenece a varios porcentajes

10.5.2.1 PORCENTAJES EXCLUYENTES

Los elementos, sujetos o respuestas sólo pueden pertenecer a una de las diferentes categorías establecidas; todas ellas sumadas suponen el cien por cien.

Las fórmulas son:

$$\frac{\text{Diferencia Observada}}{\text{Error Estándar}} = \frac{p_1 - p_2}{S_{p_1} - S_{p_2}}$$

$$S_{p_1} - S_{p_2} = \sqrt{\frac{1}{n} (p_1 q_1 + p_2 q_2 + 2 p_1 p_2)}$$

Siendo n el tamaño de la muestra y $q = 100 - p$

EJEMPLO:

Supongamos que se quiere determinar si la diferencia de porcentaje de intención de voto, obtenida con una muestra de 1.000 ciudadanos con derecho a voto, respecto al PP y PSOE, es significativa a un nivel del 5%. Los resultados obtenidos en la encuesta han sido los siguientes:

Partido	PSOE	PP	PAR	IU	RESTO
Porcentaje	32	24	16	10	18

Solución

$$p_1 = 32 \quad q_1 = 100 - 32 = 68 \quad p_2 = 24 \quad q_2 = 76$$

$$s_{p_1} - s_{p_2} = \sqrt{\frac{1}{1000} (32 * 68 + 24 * 76 + 2 * 32 * 24)} = 2'35$$

$$\frac{\text{Diferencia Observada}}{\text{Error Estándar}} = \frac{32 - 24}{2'35} = 3'4$$

Como $3'4 \geq 1'96$, la diferencia de porcentaje sí que es significativa.

10.5.2.2 PORCENTAJES SOLAPADOS

Los diferentes sujetos o respuestas pueden pertenecer a varias de las diferentes categorías establecidas, todas ellas sumadas totalizan más del 100%.

Las fórmulas son:

$$\frac{\text{Diferencia Observada}}{\text{Error Estándar}} = \frac{p_1 - p_2}{s_{p_1} - s_{p_2}}$$

$$s_{p_1} - s_{p_2} = \sqrt{\frac{1}{n} (p_1 q_1 + p_2 q_2 + 2(p_1 p_2 - p_{12}))}$$

Donde n es el tamaño de la muestra, $q = 100 - p$ y p_{12} es el porcentaje de respuesta que está a la vez en la categoría 1 y 2.

Análisis de la Investigación Cuantitativa

EJEMPLO:

En un estudio realizado sobre una muestra de 1.000 lectores de prensa dominical, se han obtenido los siguientes resultados

Prensa	A	B	C	D	E	F	G	TOTAL
Porcentaje	32	43	22	10	8	6	12	133

Si un 15% compra indistintamente el C y el D, hay que determinar si la diferencia observada es significativa a un nivel del 5%.

Solución:

$$p_1 = 22 \quad q_1 = 100 - 22 = 78 \quad p_2 = 10 \quad q_2 = 100 - 10 = 90 \quad p_{12} = 15$$

$$s_{p_1} - s_{p_2} = \sqrt{\frac{1}{1000} (22 * 78 + 10 * 90 + 2(22 * 10 - 15))} = 1'74$$

$$\frac{\text{Diferencia Observada}}{\text{Error Estándar}} = \frac{22 - 10}{1'74} = 6'9$$

Como $6'9 > 1'96$, la diferencia observada sí que es significativa a un nivel del 5%.

10.6 TEST PARA MUESTRAS RELACIONADAS

Se utiliza el test t. El cálculo del estadístico se realiza con la siguiente fórmula:

$$t = \frac{m_d - m_D}{s_d / \sqrt{n}}$$

Siendo

$$m_d = \frac{\sum_{i=1}^n d_i}{n} \quad m_D \text{ La diferencia esperada y } s_d = \sqrt{\frac{\sum_{i=1}^n (d_i - m_d)^2}{n - 1}}$$

Siendo d_i la diferencia en los valores muestrales, m_d la media de las diferencias muestrales, m_D la media de la diferencia esperada y s_d la desviación típica de las diferencias muestrales.

11. TEST NO PARAMÉTRICOS

11.1 INTRODUCCIÓN

Una prueba estadística no paramétrica está basada en un modelo que especifica sólo condiciones muy generales y ninguna acerca de la forma específica de la distribución de la población de la que fue tomada la muestra.

Los test no paramétricos tienen por objetivo el comprobar si se pueden generalizar las conclusiones obtenidas a través de una muestra al total del universo o población.

Son de aplicación cuando:

- Sólo podemos disponer de una muestra pequeña (no se puede aplicar el teorema central del límite).
- No disponemos de una medición en escala métrica. Disponemos de resultados obtenidos en escala nominal u ordinal.
- No hay exigencia de un tipo de distribución concreto.

11.2 CLASIFICACIÓN DE LOS TEST NO PARAMÉTRICOS

Se establece de acuerdo con los siguientes criterios:

1. Número de muestras que se tienen (una, dos o más).
2. Existencia o no de relación entre las muestras independientes cuando tienen varianzas muy distintas. Se considera que las muestras están relacionadas en las siguientes situaciones: cuando se entrevista a sus componentes antes y después y, asimismo cuando se entrevistan con reiteración (experimentación, paneles y, en general, en todos los estudios longitudinales). Se supone que la varianza es idéntica para los distintos momentos.
3. Escala de medida de las variables objeto de estudio (nominal u ordinal; si la variable es métrica, se integra en uno de los niveles citados anteriormente).

De acuerdo con estos criterios, la clasificación de los test no paramétricos se puede resumir en el siguiente esquema:

1 Una muestra

Escala de medida

Nominal

Ordinal

Test no paramétrico

Chi cuadrado, Rachas, Binomial

Kolmogorov – Smirnov

2 Dos muestras

2.1 Independientes

Escala de medida

Nominal

Ordinal

Test no paramétrico

Chi cuadrado

Mediana, Kolmogorov - Smirnov, Mann -
Whitney, Wald - Wolfowitz

2.2 Relacionadas

Escala de medida

Nominal

Ordinal

Test no paramétrico

McNemar

Signos, Wilcoxon

3 “h” muestras

3.1 Independientes

Escala de medida

Nominal

Ordinal

Test no paramétrico

Chi cuadrado

Mediana, Kruskal –Wallis

3.2 Relacionadas

Escala de medida

Nominal

Ordinal

Test no paramétrico

Q de Cochran

Friedman, Kendall

El principal inconveniente de los test no paramétricos es que no son tan potentes como los test paramétricos, fundamentalmente porque el nivel de exigencia en su aplicación es menor. Esto puede corregirse aumentando el tamaño de la muestra.

11.3 BREVE DESCRIPCIÓN DE DIFERENTES TEST NO PARAMÉTRICOS

11.3.1 INTRODUCCIÓN

En primer lugar explicaremos brevemente algunas pruebas estadísticas no paramétricas que se utilizan para probar una hipótesis derivada de una muestra.

Se trata de dar respuesta a las siguientes cuestiones:

- ¿Hay diferencia significativa entre la muestra y el universo o población al determinar la medida de tendencia central?
- ¿La muestra objeto de estudio fue obtenida de un universo con una forma uniforme (normal)?
- ¿Hay diferencias significativas entre las frecuencias observadas y las esperadas (en base a alguna teoría previa)?
- ¿Existe diferencia significativa entre las proporciones esperadas y las observadas en una serie de observaciones dicotómicas?
- ¿Se puede considerar que la muestra objeto de estudio corresponde a una muestra aleatoria de algún tipo de población conocida?

11.3.2 UNA MUESTRA MEDIDA UNA SOLA VEZ

En primer lugar trataremos del estudio de pruebas de bondad de ajuste para una muestra medida una sola vez. Las más utilizadas son:

1. La prueba de Chi cuadrado de una muestra
2. La prueba binomial
3. La prueba de Kolmogorov - Smirnov de una muestra

La prueba de Chi cuadrado se utiliza cuando los datos obtenidos de la muestra están en categoría discreta y cuando las frecuencias esperadas son suficientemente grandes. Cuando $k = 2$, es decir, los grados de libertad $gl = 1$, cada frecuencia esperada debe ser mayor o igual que cinco (≥ 5); y cuando k es mayor de dos ($k > 2$), no más del 20% de las frecuencias esperadas deben ser menores de cinco (5) y en ningún caso la frecuencia esperada puede ser menor de uno.

La prueba binomial es adecuada cuando hay dos categorías en la clasificación de los datos obtenidos con la muestra objeto de estudio. Es también útil cuando el tamaño de la muestra es tan pequeño que la prueba de chi cuadrado resulta inadecuada.

La prueba de Kolmogorov - Smirnov de una muestra debe emplearse cuando se puede suponer que la variable en consideración tiene una distribución continua.

11.4 TEST DE LA CHI CUADRADO

Es uno de los estadísticos más utilizados, sobre todo en las tabulaciones cruzadas; también es una medida de asociación. Se trata de una prueba de significación estadística muy adecuada para variables no métricas, es decir variables medidas en escalas nominal u ordinal.

Esta prueba consiste en comparar las frecuencias que se han obtenido en la investigación con las que desde un planteamiento teórico cabría esperar si se diera una distribución normal.

La fórmula correspondiente es:

$$\chi^2 = \sum_i \sum_j \frac{(O_{ij} - E_{ij})^2}{E_{ij}}$$

Donde

O_{ij} son las frecuencias observadas de la categoría i de una variable y de la categoría j de la otra variable

E_{ij} son las frecuencias esperadas de la categoría i de una variable y de la categoría j de la otra variable

Los requisitos para aplicar esta prueba son: las frecuencias teóricas han de tomar al menos el valor 5 en menos de un 20% de las celdas y en ningún caso la frecuencia esperada puede ser menor que 1. Cada observación debe ser independiente de las otras, no sirve con experimentos en los que se interroga antes y después del tratamiento.

En tablas de un grado de libertad o del tipo 2 x 2, se aplica la siguiente fórmula

$$\chi^2 = \frac{n \left(|ad - bc| - \frac{n}{2} \right)^2}{(a+b)(a+c)(c+d)(b+d)}$$

la tabla es del tipo:

a	b
c	d

La distribución de chi cuadrado está determinada por los grados de libertad, su media es igual al número de grados de libertad y su varianza dos veces esa cifra. Cuando los grados de libertad toman un valor alto la distribución se aproxima a la normal.

Resumen de la prueba Chi cuadrado

1. Se sitúan las frecuencias observadas dentro de k categorías
2. La suma de todas las frecuencias debe ser n (número de casos, tamaño de la muestra)
3. Partiendo de H_0 se determinan las frecuencias esperadas, teniendo en cuenta las limitaciones. Las frecuencias teóricas han de tomar al menos el valor 5 en menos de un 20% de las celdas y en ningún caso la frecuencia esperada puede ser menor que 1
4. Se determinan los grados de libertad $gl = k - n_p - 1$, donde n_p es el número de parámetros estimados de los datos y usados al calcular las frecuencias esperadas
5. Se aplica la fórmula

$$\chi^2 = \sum_i \sum_j \frac{(O_{ij} - E_{ij})^2}{E_{ij}}$$

6. Se determina la probabilidad asociada, consultando con las tablas; si ésta es igual o menor que la significación α se rechaza H_0 . Si el nivel de significación (α) es menor que el de contraste (5% ó 1%) , se rechaza la H_0

O bien si $\chi^2_{calculada} > \chi^2_{tablas}$ (se rechaza la hipótesis nula)

11.4.1 CASO PRÁCTICO

Un fabricante de refrescos quiere saber si el nuevo producto lanzado produce la misma satisfacción que el clásico entre sus consumidores. En la tabla siguiente se resumen los datos para el producto clásico y los obtenidos en una muestra de 100 usuarios del producto nuevo.

Análisis de la Investigación Cuantitativa

Se quiere conocer si las diferencias obtenidas son significativas para un nivel de significación del 5%.

	CLÁSICO	NUEVO PRODUCTO
MUY SATISFECHO	30	31
BASTANTE SATISFECHO	35	40
MODERADAMENTE SATISFECHO	25	20
ESCASAMENTE SATISFECHO	10	9

Solución

Las hipótesis de trabajo correspondientes son:

H_0 . No hay diferencia en la valoración de la satisfacción entre los consumidores del producto clásico y del nuevo

H_1 . Sí que existen diferencias

Para un nivel de significación del 5% y $4 - 1 = 3$ grados de libertad el valor en tablas para la chi cuadrado es 7'81.

Aplicando la fórmula obtenemos

$$c^2 = \sum_i \sum_j \frac{(O_{ij} - E_{ij})^2}{E_{ij}} = \frac{(31 - 30)^2}{30} + \frac{(40 - 35)^2}{35} + \frac{(20 - 25)^2}{25} + \frac{(9 - 10)^2}{10} = 1'847$$

Como el valor obtenido 1'847 es inferior al de tablas de 7'81, significa que no podemos rechazar la H_0 , luego no existen diferencias en el nivel de satisfacción entre los consumidores del producto clásico y del nuevo.

Análisis de la Investigación Cuantitativa

TABLA ESTADÍSTICA: DISTRIBUCIÓN DE χ^2
Valores de la función de distribución
g.l. = grados de libertad
 χ^2_c tal que $p(\chi^2 \leq \chi^2_c) = p$

Probabilidad p											
g.l.	0,995	0,990	0,975	0,950	0,900	0,500	0,100	0,050	0,025	0,010	0,005
1	7,88	6,63	5,02	3,84	2,71	0,45	0,01	0,00	0,00	0,00	0,00
2	10,60	9,21	7,38	5,99	4,61	1,39	0,21	0,10	0,05	0,02	0,01
3	12,84	11,34	9,35	7,81	6,25	2,37	0,58	0,35	0,22	0,12	0,07
4	14,86	13,28	11,14	9,49	7,78	3,36	1,06	0,71	0,48	0,30	0,21
5	16,75	15,09	12,83	11,17	9,24	4,25	1,61	1,15	0,83	0,55	0,41
6	18,55	16,81	14,45	12,69	10,64	5,35	2,20	1,64	1,24	0,87	0,68
7	20,28	18,48	16,01	14,07	12,02	6,35	2,83	2,17	1,69	1,24	0,99
8	21,96	20,09	17,53	15,51	13,36	7,34	3,49	2,73	2,18	1,65	1,34
9	23,59	21,67	19,02	16,92	14,68	8,34	4,17	3,33	2,70	2,09	1,73
10	25,19	23,21	20,48	18,31	15,99	9,34	4,87	3,94	3,25	2,56	2,16
11	26,76	24,73	21,92	19,68	17,28	10,34	5,58	4,57	3,82	3,05	2,60
12	28,30	26,22	23,34	21,03	18,55	11,34	6,30	5,23	4,40	3,57	3,07
13	29,82	27,69	24,74	22,36	19,81	12,34	7,04	5,89	5,01	4,11	3,57
14	31,32	29,14	26,12	23,68	21,06	13,34	7,79	6,57	5,63	4,66	4,07
15	32,80	30,58	27,49	25,00	22,31	14,34	8,55	7,26	6,26	5,23	4,60
16	34,27	32,00	28,85	26,30	23,54	15,34	9,31	7,96	6,91	5,81	5,14
17	35,72	33,41	30,19	27,59	24,77	16,34	10,09	8,67	7,56	6,41	5,70
18	37,16	34,81	31,53	28,87	25,99	17,34	10,86	9,39	8,23	7,01	6,26
19	38,58	36,29	32,85	30,14	27,20	18,34	11,65	10,12	8,91	7,63	6,84
20	40,00	37,67	34,27	31,41	28,41	19,34	12,44	10,85	9,59	8,26	7,43
21	41,40	38,93	35,48	32,67	29,62	20,34	13,24	11,59	10,28	8,90	8,03
22	42,80	40,29	36,78	33,92	30,81	21,34	14,04	12,34	10,98	9,54	8,64
23	44,18	41,64	38,08	35,17	32,01	22,34	14,85	13,09	11,69	10,20	9,26
24	45,56	42,98	39,36	36,42	33,20	23,34	15,66	13,85	12,40	10,86	9,89
25	46,93	44,31	40,65	37,65	34,38	24,34	16,47	14,61	13,12	11,52	10,52
26	48,29	45,64	41,92	38,89	35,56	25,34	17,29	15,38	13,84	12,20	11,16
27	49,64	46,96	43,29	40,11	36,74	26,34	18,11	16,15	14,57	12,83	11,81
28	50,99	48,28	44,46	41,34	37,92	27,34	18,94	16,93	15,31	13,56	12,46
29	52,34	49,59	45,72	42,56	39,09	28,34	19,77	17,71	16,05	14,26	13,12
30	53,67	50,89	46,98	43,77	40,26	29,34	20,60	18,49	16,89	14,96	13,78
40	66,77	63,69	59,34	55,76	51,81	39,34	29,05	26,51	24,43	22,16	20,71
60	91,95	88,38	83,30	79,08	74,40	59,34	46,56	43,19	40,48	37,43	35,58

11.5 PRUEBA DE LA BINOMIAL

Existe un gran número de poblaciones o universos que son binarios o dicotómicos (por ejemplo: hombre - mujer, oyente - no oyente, consumidor - no consumidor, ... etc.).

En esta situación, en el universo o población sólo existen dos categorías. Por tanto, para cada observación (x) realizada en la muestra (n) se pueden dar dos valores 1 ó 0, en función de la categoría observada.

La probabilidad de observar la primera categoría la representamos por p; por consiguiente, la probabilidad para la segunda categoría será $1 - p = q$. Esta situación la podemos representar por:

$$P(x = 1) = p \quad \text{y} \quad P(x = 0) = 1 - p = q$$

Se presupone que cada probabilidad es constante sin considerar el número de elementos observados. El valor de la proporción (π) para el universo o población es un valor fijo, sin embargo, aún conociendo el valor de éste para la población, no podemos esperar que el resultado obtenido sobre una muestra aleatoria (p) coincida exactamente con el valor del de la población (π).

La distribución binomial se utiliza para determinar las probabilidades de los resultados obtenidos al estudiar una muestra procedente de una población dicotómica.

Metodología:

Se establece la hipótesis nula como: $H_0 : p = \pi$

La prueba nos dirá si es razonable creer que las proporciones (frecuencias) de las categorías obtenidas en una muestra (n) han sido extraídas de una muestra correspondiente a una población con valores hipotéticos π y $1 - \pi$.

Si consideramos los resultados de la distribución binomial como 1 para el éxito y 0 para el fracaso, el número de éxitos vendrá dado por:

$$Y = \sum_{i=1}^n x_i$$

En una muestra de tamaño n, la probabilidad de obtener k elementos de una categoría y n - k de la otra será:

$$P(Y = k) = \binom{n}{k} p^k q^{n-k}$$

Siendo: n el tamaño de la muestra, $k = 1, 2, 3, \dots, n$, p la proporción de observaciones

para $x = 1$, q la proporción de observaciones para $x = 0$ y $\binom{n}{k} = \frac{n!}{k!(n-k)!}$

EJEMPLO:

Supongamos que lanzamos un dado 5 veces. ¿Cuál es la probabilidad de que dos de las tiradas sea un seis?

En esta situación $n = 5$, $k = 2$, (número de observaciones que corresponden al seis),

$p = \frac{1}{6}$ y $q = \frac{5}{6}$ la variable aleatoria $Y = k = 2$, aplicando la fórmula anterior obtenemos:

$$P(Y = k) = \binom{n}{k} p^k q^{n-k} = \frac{1 \times 2 \times 3 \times 4 \times 5}{(1 \times 2)(1 \times 2 \times 3)} \left(\frac{1}{6}\right)^2 \left(\frac{5}{6}\right)^3 = 0.16$$

En la praxis, cuando probamos una hipótesis la cuestión no suele ser ¿cuál es la probabilidad de obtener exactamente los valores observados?, sino que la pregunta es del tipo siguiente:

¿cuál es la probabilidad de obtener valores tan extremos o más extremos que los valores observados?

La probabilidad deseada en este caso es:

$$P(Y \geq k) = \sum_{i=k}^n \binom{n}{i} p^i q^{n-i}$$

Es decir, sumamos la probabilidad de los resultados observados con la probabilidad de resultados más extremos. Siguiendo con el ejemplo, el planteamiento de la cuestión es: Determinar la probabilidad de obtener dos o menos seises cuando hacemos cinco lanzamientos con un dado normal.

Esto significa que deberemos obtener la probabilidad de sacar 0, 1 y 2 seises, aplicando la fórmula anterior, y recordando que por definición $0! = 1$ y $x^0 = 1$, obtenemos

$P(Y \leq 2) = P(Y = 0) + P(Y = 1) + P(Y = 2)$ sustituyendo obtenemos

$$P(Y = 0) = \frac{5!}{0!5!} \left(\frac{1}{6}\right)^0 \left(\frac{5}{6}\right)^5 = 0'40 \quad P(Y = 1) = \frac{5!}{1!4!} \left(\frac{1}{6}\right)^1 \left(\frac{5}{6}\right)^4 = 0'40 \text{ y}$$
$$P(Y = 2) = \frac{5!}{2!3!} \left(\frac{1}{6}\right)^2 \left(\frac{5}{6}\right)^3 = 0'16$$

Luego $P(Y \leq 2) = 0'40 + 0'40 + 0'16 = 0'96$

11.6 PRUEBA BINOMIAL PARA MUESTRAS PEQUEÑAS

Cuando en la Investigación Comercial se trabaja con variables dicotómicas es muy frecuente utilizar como hipótesis nula $H_0 : p = 1/2$. Cuando tenemos muestras pequeñas, es decir $n \leq 30$, se utilizan tablas ya calculadas al efecto considerando $H_0 : p = 1/2$. Este tipo de tablas nos indican las probabilidades asociadas con la ocurrencia de diferentes valores tan pequeños como k para diferentes valores de n .

Las probabilidades proporcionadas en las tablas (ver tabla) son unidireccionales. Se utiliza la prueba unidireccional cuando se predice con anterioridad cual de las dos categorías (1 y 0) contendrá el número más pequeño de casos (k). Cuando la predicción es simplemente que las dos frecuencias difieran, se utiliza la prueba bidireccional; en este caso, los valores de la tabla se duplican.

Debido a la simetría de la distribución binomial, cuando $p = 1/2$ se cumple:

$$P(Y \geq k) = P(Y \leq n - k)$$

11.7 PRUEBA BINOMIAL PARA MUESTRAS GRANDES

Una muestra se considera grande cuando su tamaño es mayor de 30 elementos¹⁰ ($n > 30$). Cuando se incrementa el tamaño de la muestra la distribución binomial tiende a convertirse en la distribución normal. Es decir, al aumentar el tamaño de la muestra n la distribución de la variable Y se aproxima a la distribución normal. La tendencia es rápida cuando $p = 1/2$, y lenta conforme más se aproxima al valor 1 ó 0. Como aproximación podemos usar la siguiente :

Regla estadística: Si $npq > 9$ la prueba estadística basada en la normal es suficientemente exacta para ser usada.

¹⁰Algunos autores consideran 35 elementos.

Con las consiguientes limitaciones la distribución muestral de Y es aproximadamente normal con media $m = np$ y varianza $s^2 = npq$. El estadístico usado en este caso será

$$Z_Y = \frac{y - m}{s} = \frac{y - np}{\sqrt{npq}}$$

Esta aproximación a la distribución normal mejora si se usa una “corrección por continuidad”. Esta corrección es necesaria debido a que la distribución normal es continua, mientras que la distribución binomial corresponde a variables discretas.

Para realizar esta corrección o ajuste se considera la frecuencia observada Y ocupando un intervalo, cuyos límites inferior y superior se encuentran media unidad por debajo o por encima, respectivamente, de la frecuencia observada. Por consiguiente si $Y < np$ agregaremos 0.5 y si $Y > np$ restaremos 0.5. Por tanto la diferencia observada es reducida por 0.5.

La fórmula correspondiente a la razón crítica una vez realizado este ajuste quedará como sigue:

$$Z_Y = \frac{(y \pm 0.5) - np}{\sqrt{npq}}$$

Potencia eficacia

Debido a que para variables dicotómicas no existe una prueba paramétrica aplicable, no tiene sentido el hablar de la potencia eficacia.

Si estudiamos una variable continua que es dicotomizada y se utiliza la prueba binomial con los datos resultantes, la prueba puede perder información. En este caso, la prueba tiene una potencia eficacia del 95% para $n = 6$ disminuyendo al aumentar la muestra hasta producirse una eficacia asintótica del 63% ($2/\pi$).

Resumen de la prueba binomial

Los pasos para la utilización de la prueba binomial considerando $H_0 : p = \frac{1}{2}$ son:

1. Determinar el número de casos estudiados (n)
2. Determinar la frecuencia de cada una de las dos categorías
3. En función del tamaño de la muestra (muestra pequeña o grande), determinar el valor en tablas en función del nivel de significación α
4. Si la probabilidad asociada con el valor observado de Y, o valores aún más extremos, es igual o menor que el correspondiente valor α se rechaza H_0 , en caso contrario no se rechaza

Los supuestos básicos de esta prueba son la independencia de las observaciones y que la probabilidad permanece constante durante el estudio.

CASO PRÁCTICO

El responsable de una agencia de publicidad encargada de una campaña de imagen de un Gobierno Autónomo asegura haber contactado con el 60% de la población. Para verificarlo se realiza un estudio sobre una muestra de 25 personas, a quienes se pregunta si conocen la campaña. Cuatro de estas personas declaran conocerla. ¿Se puede decir que la proporción del 60% ha sido contactada a un nivel del 55?

La probabilidad del resultado es

$$P(y \leq 4) = P(0) + P(1) + P(2) + P(3) + P(4)$$

Calculamos las diferentes probabilidades aplicando la fórmula

$$P(Y = k) = \binom{n}{k} p^k q^{n-k}$$

Sustituyendo obtenemos:

$$\begin{aligned} P(0) &= \frac{25!}{0!(25-0)!} 0'6^0 0'4^{25} & P(1) &= \frac{25!}{1!(25-1)!} 0'6^1 0'4^{24} \\ P(2) &= \frac{25!}{2!(25-2)!} 0'6^2 0'4^{23} & P(3) &= \frac{25!}{3!(25-3)!} 0'6^3 0'4^{22} \\ P(4) &= \frac{25!}{4!(25-4)!} 0'6^4 0'4^{21} \end{aligned}$$

Realizando los correspondientes cálculos obtenemos

$$P(Y = 4) = 0'000008165$$

Análisis de la Investigación Cuantitativa

Los resultados obtenidos se recogen en la siguiente tabla:

Tabla de resultados

N=Exitos	N!	P	Q	Intentos-N	P(N)	P(N) acum
0	1	0,6	0,4	25	0,000000000	0,000000000
1	1			24	0,000000004	0,000000004
2	2			23	0,000000076	0,000000080
3	6			22	0,000000874	0,000000954
4	24			21	0,000007210	0,000008165
5	120			20	0,000045425	0,000053590
6	720			19	0,000227126	0,000280715
7	5040			18	0,000924725	0,001205441
8	40320			17	0,003120948	0,004326388
9	362880			16	0,008842685	0,013169073
10	3628800			15	0,021222445	0,034391518
11	39916800			14	0,043409546	0,077801064
12	479001600			13	0,075966705	0,153767769
13	6227020800			12	0,113950058	0,267717827
14	87178291200			11	0,146507217	0,414225044
15	1307674368000			10	0,161157939	0,575382982
16	20922789888000			9	0,151085568	0,726468550
17	355687428096000			8	0,119979715	0,846448265
18	6402373705728000			7	0,079986477	0,926434742
19	121645100408832000			6	0,044203053	0,970637795
20	2432902008176640000			5	0,019891374	0,990529169
21	51090942171709400000			4	0,007104062	0,997633231
22	1124000727777610000000			3	0,001937471	0,999570703
23	25852016738885000000000			2	0,000379071	0,999949773
24	620448401733239000000000			1	0,000047384	0,999997157
25	15511210043331000000000000			0	0,000002843	1,000000000

Si comparamos con la tabla de la distribución binomial para $N = 25$, obtenemos la siguiente

Conclusión

Si analizamos el resultado con el obtenido en las tablas que es:

para $p = 0,55$ $P(Y = 4) = 0,000000001$ como la calculada es $P(Y = 4) = \mathbf{0,000008165}$

Como la p calculada es mayor que la de tablas NO podemos rechazar la hipótesis nula (H_0)

Análisis de la Investigación Cuantitativa

TABLA DE LA DISTRIBUCIÓN BINOMIAL PARA N = 25

En esta tabla se han omitido los decimales. Las entradas deben leerse como 0'0000. Para valores $p > 0.5$ se usa la parte inferior para p y la columna derecha para k

0	7778	2774	718	172	38	8	1	0	0	0	0	25	25
1	1964	3650	1994	759	236	63	14	5	0	0	0	24	
2	238	2305	2659	1607	708	251	74	30	4	1	0	23	
3	18	930	2265	2174	1358	641	243	114	19	4	1	22	
4	1	269	1384	2110	1867	1175	572	313	71	18	4	21	
5	0	60	646	1564	1960	1645	1030	658	199	63	16	20	
6	0	10	239	920	1633	1828	1472	1096	442	172	53	19	
7	0	1	72	441	1108	1654	1712	1487	800	381	143	18	
8	0	0	18	175	623	1241	1651	1673	1200	701	322	17	
9	0	0	4	58	294	781	1336	1580	1511	1084	609	16	
10	0	0	1	16	118	417	916	1264	1612	1419	974	15	
11	0	0	0	4	40	189	536	862	1465	1583	1328	14	
12	0	0	0	1	12	74	268	503	1140	1511	1550	13	
13	0	0	0	0	3	25	115	251	760	1236	1550	12	
14	0	0	0	0	1	7	42	108	434	867	1328	11	
15	0	0	0	0	0	2	13	40	212	520	974	10	
16	0	0	0	0	0	0	4	12	88	266	609	9	
17	0	0	0	0	0	0	1	3	31	115	322	8	
18	0	0	0	0	0	0	0	1	9	42	143	7	
19	0	0	0	0	0	0	0	0	2	13	53	6	
20	0	0	0	0	0	0	0	0	0	3	16	5	
21	0	0	0	0	0	0	0	0	0	1	4	4	
22	0	0	0	0	0	0	0	0	0	0	1	3	
23	0	0	0	0	0	0	0	0	0	0	0	2	
24	0	0	0	0	0	0	0	0	0	0	0	1	
25	0	0	0	0	0	0	0	0	0	0	0	0	
	0.99	0.95	0.90	0.85	0.80	0.75	0.70	2/3	0.60	0.55	0.50	k	N
	D												

11.8 TEST DE KOLMOGOROV – SMIRNOV (KS)

Este tipo de prueba trata de ver el grado de acuerdo entre la distribución de un conjunto de valores obtenidos a través de una muestra y alguna distribución teórica específica. Para utilizar esta prueba las variables deben estar medidas al menos en una escala ordinal.

Es una prueba parecida a la de Chi cuadrado, que consiste en comparar valores observados de una variable, con valores esperados calculados a priori.

La H_0 es la ausencia de diferencias entre los valores observados y los esperados.

Metodología:

La prueba KS supone que la distribución de las variables que van a ser probadas es continua (está especificada por la distribución de frecuencias acumuladas).

Para su desarrollo se procede como sigue:

- Se disponen ordenadamente las frecuencias observadas y las esperadas
- Se calculan las frecuencias relativas acumuladas, tanto las observadas como las esperadas
- Se determinan las diferencias entre las frecuencias relativas acumuladas
- Se toma la mayor diferencia en términos absolutos, es decir, el valor que hace máximo:

$$D = \max. | O_i - E_i |$$

- Se fija el nivel de significación α
- Se compara el valor obtenido D con el valor de tablas KS para el nivel de significación elegido y muestra de tamaño inferior a 35 elementos.

Si la muestra es grande, es decir tiene más de 35 elementos, la H_0 se rechaza de acuerdo con los siguientes criterios:

$$\text{Para } \alpha = 0'10 \quad D \geq \frac{1'22}{\sqrt{n}} \quad \text{Para } \alpha = 0'05 \quad D \geq \frac{1'36}{\sqrt{n}} \quad \text{y}$$

$$\text{Para } \alpha = 0'01 \quad D \geq \frac{1'63}{\sqrt{n}}$$

Potencia eficacia

La prueba KS de una muestra trata las observaciones de forma individual (por separado) y por ello no necesariamente pierde información al hacer la combinación de categorías, aunque puede ser conveniente usar agrupaciones de variables.

Cuando trabajamos con muestras pequeñas, es una prueba exacta, mientras que la de Chi cuadrado es sólo aproximada. Para muestras grande, ambas pruebas, KS y Chi cuadrado, dan resultados similares.

Resumen de la prueba KS

- La distribución teórica se especifica según H_0
- Las frecuencias observadas y las teóricas se convierten en frecuencias relativas acumuladas
- Se aplica $D = \max. |O_i - E_i|$
- Comparamos con tablas. Buscamos la probabilidad asociada (bidireccional, dos colas) con la ocurrencia según H_0 . Si esta probabilidad es igual o menor que α , se rechaza H_0

CASO PRÁCTICO

Un fabricante va a lanzar un nuevo producto con un alto componente ecológico. Quiere conocer la importancia que dan los consumidores al componente ecológico en su decisión de compra, por lo que realiza un estudio sobre una muestra de 100 consumidores potenciales.

La pregunta a los consumidores, acerca de la cuestión, se realiza mediante una escala de cinco puntos (5), siendo 1 ninguna importancia y 5 mucha importancia. Queremos conocer si hay diferencia significativa entre los valores observados y los esperados para un nivel de significación del 5%.

Los resultados se resumen en el siguiente cuadro:

Análisis de la Investigación Cuantitativa

Resultados de la encuesta

Categoría	Frecuencia	Observada		Esperada		D
		%	% acumul	%	% acumul	
1	8	0'08	0'08	0'20	0'20	0'12
2	15	0'15	0'23	0'20	0'40	0'17
3	17	0'17	0'40	0'20	0'60	0'20
4	35	0'35	0'75	0'20	0'80	0'05
5	25	0'25	1'00	0'20	1'00	0'00
Total	100	1,00		1'00		

Solución:

Si el efecto ecológico no existiera en la decisión de compra por parte del consumidor, ésta se repartiría por igual entre las diferentes categorías de respuesta, siendo $100 : 5 = 20\%$ para cada categoría.

La hipótesis nula será H_0 : no existe diferencia entre los valores observados y los esperados.

Aplicamos la prueba KS para un nivel de significación del 5%. El valor crítico viene dado por:

$$D = \frac{1'36}{\sqrt{n}} = \frac{1'36}{\sqrt{100}} = 0'136$$

El valor crítico observado viene dado por:

$$D = \max_i |O_i - E_i| = \max |0'4 - 0'6| = 0'2$$

Conclusión:

Como el valor observado 0'20 es mayor que el teórico 0'136, se rechaza la hipótesis nula. Es decir, el componente ecológico sí que tiene importancia en la decisión de compra de este producto por parte de los consumidores.

Análisis de la Investigación Cuantitativa

TABLA DE KOLMOGOROV – SMIRNOV

N	0.2	0.15	0.10	0.05	0.01
1	.900	.925	.950	.975	.995
2	.684	.726	.776	.842	.929
3	.565	.597	.642	.708	.828
4	.494	.525	.564	.624	.733
5	.446	.474	.510	.565	.669
6	.410	.436	.470	.521	.618
7	.381	.405	.438	.486	.577
8	.358	.381	.411	.457	.543
9	.339	.360	.388	.432	.514
10	.322	.342	.368	.410	.490
11	.307	.326	.352	.391	.468
12	.295	.313	.338	.375	.450
13	.284	.302	.325	.361	.433
14	.274	.292	.314	.349	.418
15	.266	.283	.304	.338	.404
16	.258	.274	.295	.328	.392
17	.250	.266	.286	.318	.381
18	.244	.259	.278	.309	.371
19	.237	.252	.272	.301	.363
20	.231	.246	.264	.291	.356
25	.21	.22	.24	.27	.32
30	.19	.20	.22	.24	.29
35	.18	.19	.21	.23	.27
Mas de 35	$1.07/\sqrt{n}$	$1.14/\sqrt{n}$	$1.22/\sqrt{n}$	$1.36/\sqrt{n}$	$1.63/\sqrt{n}$

11.9 CASO DE UNA MUESTRA MEDIDA DOS VECES

Se trata de pruebas que se utilizan para situaciones de prueba antes y después. En este tipo de estudios se mide al mismo individuo en ocasiones sucesivas, actuando el mismo individuo como control. Las más usuales son:

- Test de McNemar
- Prueba de los signos
- Prueba de rangos asignados de Wilcoxon

Seguidamente vamos a ver en qué consisten.

11.9.1 TEST DE MCNEMAR

Parte de una situación dicotómica en la que se aplica un tratamiento y se registra la nueva situación para comprobar los cambios producidos. La medición es, como mínimo, en escala nominal. Este tipo de prueba es interesante en estudios panel y en experimentación.

Metodología:

Para probar la significación de cualquier cambio observado se utiliza una tabla 2 x 2, que representa el primer y segundo grupo de respuestas de los mismos individuos. La tabla es del tipo siguiente:

DESPUÉS			
ANTES		NEGATIVO (-)	POSITIVO (+)
	POSITIVO (+)	A	B
	NEGATIVO (-)	C	D

A es el número de respuestas que fueron positivas en la primera medición y negativas en la segunda ocasión

B es la frecuencia de individuos que respondieron en positivo en las dos ocasiones

C es la frecuencia de individuos que respondieron en negativo en las dos ocasiones

D es el número de respuestas que fueron negativas en la primera medición y positivas en la segunda ocasión

$A + D$ es el total de individuos que cambiaron de respuesta. Si esta suma es menor de 10, se utiliza la prueba binomial.

La hipótesis nula es la no existencia de diferencias. H_0 que el número de cambios en cada dirección es el mismo. Es decir,

$$\frac{A+D}{2} \text{ cambiaron de } + \text{ a } - \text{ y } \frac{A+D}{2} \text{ cambiaron de } - \text{ a } +$$

Esto significa que si H_0 es verdadera, la frecuencia esperada en cada una de las celdas

será: $\frac{A+D}{2}$

Esto significa que la distribución obtenida lo hace como una χ^2 con un grado de libertad.

Si aplicamos el correspondiente estadístico y sustituimos por los valores de la tabla obtenemos:

$$c^2 = \sum_{i=1}^k \frac{(O_i - E_i)^2}{E_i} = \frac{\left(A - \frac{A+D}{2}\right)^2}{\frac{A+D}{2}} + \frac{\left(D - \frac{A+D}{2}\right)^2}{\frac{A+D}{2}} = \frac{(A-D)^2}{A+D}$$

Esta fórmula se hace más precisa cuando se efectúa la corrección por continuidad. La fórmula correspondiente es:

$$c^2 = \frac{[|A-D|-1]^2}{A+D}$$

Si el valor calculado de la Chi cuadrado es igual o mayor que el de tablas, se rechaza H_0

La prueba de McNemar se aproxima a la distribución Chi cuadrado sólo cuando el tamaño de la muestra es grande.

Potencia eficacia

No tiene sentido hablar de potencia cuando se utiliza con variables en escala nominal ya que no hay alternativas con las que comparar la prueba. Cuando las medidas y otros

Análisis de la Investigación Cuantitativa

aspectos de los datos son tales que es posible aplicar la prueba t, la prueba de McNemar tiene una eficacia del 95% para $A + D = 1$; conforme disminuye $A + D$ la potencia eficacia va decreciendo, volviéndose asintótica al 63%.

Resumen de la prueba de McNemar

- Se colocan las frecuencias observadas en una tabla 2 x 2
- Se calcula $A + D$ (si es menor de 10, se utilizará la prueba binomial)
- Si $A + D > 10$, se calcula Chi cuadrado para 1 grado de libertad, aplicando la fórmula:

$$\bullet \quad \chi^2 = \frac{[|A - D| - 1]^2}{A + D}$$

•

- El resultado obtenido se compara con tablas. Si utilizamos la prueba de una sola cola, se divide por dos el valor obtenido en tablas. Si el valor de la probabilidad de la tabla para el valor observado con $gl = 1$ es menor o igual que el asociado para H_0 , se rechaza ésta.

CASO PRÁCTICO

Se quiere conocer la intención de voto de una población en relación a un candidato determinado. Para ello, se solicita ésta intención a una muestra formada por 100 ciudadanos con derecho a voto, obteniéndose los siguientes resultados:

Intención de voto	Frecuencia
Le votaría (+)	40
No le votaría (-)	60
Total	100

Después de producirse un debate televisivo en el que participa el candidato, se pregunta a la misma muestra, obteniéndose los siguientes resultados:

Análisis de la Investigación Cuantitativa

Antes del debate	Después del debate
Votaría (+) 40	Votaría (+) 38 No votaría (-) 2
No votaría (-) 60	Votaría (+) 20 No votaría (-) 40
Total 100	Total 100

Con los resultados obtenidos realizamos la siguiente tabla de conclusiones:

		Después	
		Negativo(-)	Positivo (+)
Antes	Positivo (+)	2	38
	Negativo (-)	40	20

Se quiere conocer si el resultado obtenido es significativo para $\alpha = 5\%$

Solución:

La hipótesis nula correspondiente es: H_0 no hay diferencia en la intención de voto antes y después del debate televisivo.

Aplicamos la prueba de McNemar, cuyo estadístico es;

$$c^2 = \frac{[|A - D| - 1]^2}{A + D}$$

Donde $A = 2$ y $D = 20$. Obtenemos que

$$c^2 = \frac{[|A - D| - 1]^2}{A + D} = \frac{[|2 - 20| - 1]^2}{2 + 20} = \frac{17^2}{22} = 13.136$$

El valor en tablas es: χ^2 para $\alpha = 5\%$ y 1 grado de libertad es: 3.84

Conclusión:

Como el valor de Chi cuadrado observado es mayor que el valor teórico, se rechaza la hipótesis nula. La diferencia de intención de voto después del debate televisivo sí que es significativa al nivel del 5%.

Análisis de la Investigación Cuantitativa

TABLA ESTADÍSTICA: DISTRIBUCIÓN DE χ^2

Valores de la función de distribución

g.l. = grados de libertad

χ^2_c tal que $p(\chi^2 \leq \chi^2_c) = p$

g.l.	Probabilidad p										
	0,995	0,990	0,975	0,950	0,900	0,500	0,100	0,050	0,025	0,010	0,005
1	7,88	6,63	5,02	3,84	2,71	0,45	0,01	0,00	0,00	0,00	0,00
2	10,60	9,21	7,38	5,99	4,61	1,39	0,21	0,10	0,05	0,02	0,01
3	12,84	11,34	9,35	7,81	6,25	2,37	0,58	0,35	0,22	0,12	0,07
4	14,86	13,28	11,14	9,49	7,78	3,36	1,06	0,71	0,48	0,30	0,21
5	16,75	15,09	12,83	11,17	9,24	4,25	1,61	1,15	0,83	0,55	0,41
6	18,55	16,81	14,45	12,69	10,64	5,35	2,20	1,64	1,24	0,87	0,68
7	20,28	18,48	16,01	14,07	12,02	6,35	2,83	2,17	1,69	1,24	0,99
8	21,96	20,09	17,53	15,51	13,36	7,34	3,49	2,73	2,18	1,65	1,34
9	23,59	21,67	19,02	16,92	14,68	8,34	4,17	3,33	2,70	2,09	1,73
10	25,19	23,21	20,48	18,31	15,99	9,34	4,87	3,94	3,25	2,56	2,16
11	26,76	24,73	21,92	19,68	17,28	10,34	5,58	4,57	3,82	3,05	2,60
12	28,30	26,22	23,34	21,03	18,55	11,34	6,30	5,23	4,40	3,57	3,07
13	29,82	27,69	24,74	22,36	19,81	12,34	7,04	5,89	5,01	4,11	3,57
14	31,32	29,14	26,12	23,68	21,06	13,34	7,79	6,57	5,63	4,66	4,07
15	32,80	30,58	27,49	25,00	22,31	14,34	8,55	7,26	6,26	5,23	4,60
16	34,27	32,00	28,85	26,30	23,54	15,34	9,31	7,96	6,91	5,81	5,14
17	35,72	33,41	30,19	27,59	24,77	16,34	10,09	8,67	7,56	6,41	5,70
18	37,16	34,81	31,53	28,87	25,99	17,34	10,86	9,39	8,23	7,01	6,26
19	38,58	36,29	32,85	30,14	27,20	18,34	11,65	10,12	8,91	7,63	6,84
20	40,00	37,67	34,27	31,41	28,41	19,34	12,44	10,85	9,59	8,26	7,43
21	41,40	38,93	35,48	32,67	29,62	20,34	13,24	11,59	10,28	8,90	8,03
22	42,80	40,29	36,78	33,92	30,81	21,34	14,04	12,34	10,98	9,54	8,64
23	44,18	41,64	38,08	35,17	32,01	22,34	14,85	13,09	11,69	10,20	9,26
24	45,56	42,98	39,36	36,42	33,20	23,34	15,66	13,85	12,40	10,86	9,89
25	46,93	44,31	40,65	37,65	34,38	24,34	16,47	14,61	13,12	11,52	10,52
26	48,29	45,64	41,92	38,89	35,56	25,34	17,29	15,38	13,84	12,20	11,16
27	49,64	46,96	43,29	40,11	36,74	26,34	18,11	16,15	14,57	12,83	11,81
28	50,99	48,28	44,46	41,34	37,92	27,34	18,94	16,93	15,31	13,56	12,46
29	52,34	49,59	45,72	42,56	39,09	28,34	19,77	17,71	16,05	14,26	13,12
30	53,67	50,89	46,98	43,77	40,26	29,34	20,60	18,49	16,89	14,96	13,78
40	66,77	63,69	59,34	55,76	51,81	39,34	29,05	26,51	24,43	22,16	20,71
60	91,95	88,38	83,30	79,08	74,40	59,34	46,56	43,19	40,48	37,43	35,58

11.9.2 TEST DE LOS SIGNOS

La aplicación de esta prueba requiere un nivel de medida de, al menos, la escala ordinal. Este test se basa en la dirección de las diferencias entre dos mediciones (en este caso, de una muestra “antes-después”). Es de aplicación en investigaciones donde las mediciones cuantitativas son imposibles de realizar o no son viables. Lo que sí puede determinar para cada par de observaciones es cuál es mayor en algún sentido.

Se constatan las diferencias entre los signos positivos y negativos, si las hay, y si éstas pueden ser debidas al azar. Las variables han de ser independientes, la escala de medida como mínimo ordinal y se presupone que la variable objeto de estudio tiene una distribución continua.

La prueba no hace suposiciones acerca de la forma de la distribución, ni tampoco supone que los elementos pertenezcan al mismo universo.

Metodología:

La hipótesis nula es la ausencia de cambios entre antes y después. La podemos representar por $H_0 : P | x_i > y_i | = P | x_i < y_i | = 1/2$

Donde x_i e y_i son las respuestas obtenidas de la muestra en el momento 1 y en el momento 2, respectivamente.

La H_0 se puede plantear de la siguiente forma: la mediana de las diferencias entre x e y es cero.

Se obtienen los valores antes y después

Se determina la diferencia entre estos valores (d)

Entre el número de diferencia $d(+)$ y $d(-)$, se toma el menor.

La probabilidad de que en **n (número de diferencias producidas)** ocasiones se obtengan diferencias se compara con la binomial para $p = q = 0.5$, de forma que si es igual o menor que un valor predeterminado de α (normalmente 0.05), la diferencia es significativa, es decir, se rechaza la H_0 . Si la diferencia es mayor, no se rechaza la hipótesis nula. También se rechaza la H_0 si ocurren pocas diferencias con el mismo signo.

11.9.2.1 APLICACIÓN EN MUESTRAS PEQUEÑAS.

La probabilidad de ocurrencia de un número de positivos (+) y de negativos (-), puede determinarse recurriendo a la distribución binomial con $p = q = 0'5$, siendo n el número de pares.

Si aparecen pares en los que no hay diferencias, esto es, si no existen signos, estos datos son excluidos, reduciéndose el tamaño de la muestra. Algunos autores definen esta situación como de empates.

Las tablas nos proporcionan las probabilidades asociadas de ocurrencia del suceso, de acuerdo con los valores tan pequeños como d para $n \leq 30$ (Algunos proponen 35).

La prueba de los signos puede ser tanto unidireccional como bidireccional. Para este último caso los valores de la probabilidad deben duplicarse.

11.9.2.2 APLICACIÓN EN MUESTRAS GRANDES

En este caso n (**número de diferencias producidas**) es mayor de 30. En esta situación, se suele utilizar la aproximación normal a la distribución binomial.

La distribución tiene

$$media = m_d = np = \frac{n}{2} \text{ y varianza} = s_d^2 = npq = \frac{n}{4}$$

El valor de la razón crítica será:

$$Z = \frac{x - m}{s} = \frac{d - \frac{n}{2}}{\frac{\sqrt{n}}{2}} = \frac{2d - n}{\sqrt{n}}$$

Esta fórmula se corrige por continuidad quedando:

$$Z = \frac{(d \pm 0'5) - \frac{n}{2}}{\frac{\sqrt{n}}{2}} = \frac{2d \pm 1 - n}{\sqrt{n}}$$

Se utiliza $d + 0'5$ cuando $d < n / 2$, y $d - 0'5$ cuando $d > n / 2$

El valor obtenido de Z se considera una distribución normal, con media 0 y varianza 1.

La significación se determina comparando el valor calculado con el de tablas

Potencia eficacia

La potencia eficacia de esta prueba para $n = 6$ es del 95%, disminuyendo al aumentar n hasta hacerse asintótica al 63%.

Resumen prueba de los signos

- Se determina el signo de la diferencia para cada par
- Se calcula el valor de n . Los empates se excluyen del análisis
- El método para determinar la probabilidad de ocurrencia cuando H_0 es verdadera depende del tamaño de n
- Si $n < 30$ se utiliza la tabla de la binomial, que indica la probabilidad asociada (una cola) con valores tan pequeños observados de d (número menor de signos), para una región de rechazo de dos colas se duplica la probabilidad proporcionada por la tabla
- Si $n > 30$ se utiliza el valor z de la distribución normal. La tabla nos muestra la probabilidad asociada (unidireccional) a los valores de z . En el caso bidireccional, se duplica la probabilidad obtenida por la tabla.
- Si la probabilidad mostrada por la prueba es menor o igual a α (normalmente 0'05), se rechaza la hipótesis nula (H_0).

CASO PRÁCTICO 1

En una reunión con 25 delegados sindicales se pide la opinión de todos acerca de las nuevas medidas de seguridad en el trabajo. Sus opiniones se recoge a través de una escala de cinco puntos siendo 1 “muy desfavorable” y 5 “muy favorable”. Los delegados reciben un curso de formación después del cual, se vuelve a pedir su opinión. Los resultados obtenidos se recogen en la siguiente tabla.

Se quiere conocer si la diferencia observada es significativa para $\alpha = 0'05$.

Solución

Del análisis de la tabla obtenemos: $d(+) = 8$, $d(-) = 4$. Sin diferencia 13. El tamaño de $n = 12$

Seleccionamos la menor, esto es $d(-) = 4$. La probabilidad de obtener 4 cambios en 12 ocasiones, la buscamos en tablas y obtenemos que es: 0'194; para la prueba

Análisis de la Investigación Cuantitativa

bidireccional será 0'388. Este valor es superior al nivel de significación 0'05 que hemos seleccionado, luego no rechazamos la hipótesis nula.

Conclusión: no se producen cambios significativos en la opinión de los delegados sindicales después de recibir el curso de formación.

Tabla de resultados

Persona	Momento 1	Momento 2	Signo (d)
1	3	3	0
2	4	4	0
3	2	2	0
4	3	2	-
5	3	4	+
6	2	2	0
7	2	2	0
8	5	5	0
9	5	5	0
10	3	3	0
11	2	3	+
12	5	4	-
13	2	4	+
14	5	5	0
15	3	2	-
16	2	4	+
17	4	4	0
18	3	4	+
19	3	4	+
20	2	4	+
21	4	3	-
22	3	4	+
23	2	2	0
24	3	3	0
25	5	5	0

CASO PRÁCTICO 2

Supongamos que repetimos la experiencia anterior con una muestra de 150 elementos. Los resultados obtenidos quedan resumidos en el siguiente cuadro:

No cambian	56
Cambio positivo	32
Cambio negativo	62

Al tratarse de una muestra grande, $n = 94$, aplicamos el estadístico Z , cuya fórmula para este caso es:

$$Z = \frac{2d \pm 1 - n}{\sqrt{n}}$$

Sustituyendo obtenemos

$$Z = \frac{2d \pm 1 - n}{\sqrt{n}} = \frac{2 \times 32 + 1 - 94}{\sqrt{94}} = \frac{-29}{9'695} = -2'991$$

La probabilidad asociada al valor obtenido es: 0'0014; para dos colas sería 0'0028.

Como este valor es más pequeño que $\alpha = 0'05$, la decisión es rechazar la hipótesis nula.

Conclusión: se producen cambios significativos en la opinión de los delegados sindicales después de recibir el curso de formación.

11.9.3 TEST DE RANGOS ASIGNADOS DE WILCOXON

Se trata de una prueba parecida a la anterior, con la diferencia de que en este test se adjudica más peso a los pares que muestran mayores diferencias entre las dos condiciones que a los pares cuya diferencia es menor.

Esta prueba es de utilidad cuando se trata de emitir juicios del tipo “mayor que”. Con esta prueba el investigador puede:

Determinar qué miembro del par es mayor que

Establecer rangos en las diferencias en orden de tamaño absoluto

Se llega a considerar a las diferencias como si se correspondiesen con una medida de intervalo (en realidad son diferencias de rangos).

Investigación Comercial

Metodología

La hipótesis nula es que la suma de rangos es nula. Esto es, que la diferencia entre un sentido y otro es la misma.

Los tratamientos efectuados en el par se denominan X e Y, la diferencia de resultados por d, esto es $d_i = x_i - y_i$

Se calculan todas las diferencias

Se ponen todas las diferencias en columna sin tener en cuenta el signo. Se adjudica el rango 1 a la más pequeña. No se tiene en cuenta el signo. Se trabaja en valores absolutos (en las diferencias negativas la más pequeña será -1).

El rango se pone de acuerdo con los valores absolutos de d_i ; luego aplicaremos el signo en función de que el valor de la distancia sea positivo o negativo. De esta forma se pueden identificar los rangos de las diferencias negativas de los rangos de las diferencias positivas e indicarlos.

Tal y como decíamos anteriormente, la hipótesis nula es que los tratamientos X e Y son equivalentes, es decir, tienen la misma mediana y la misma distribución continua.

Si H_0 es verdadera, la suma de los rangos de signo positivo será la misma que la suma de rangos de tipo negativo. Por consiguiente, si la diferencia es muy distinta podemos deducir que el tratamiento X difiere del Y, y por tanto, rechazaríamos H_0 .

Rechazaremos la H_0 siempre que la suma de los rangos positivos o negativos sea muy pequeña.

Los estadísticos que se utilizan en esta prueba son:

T^+ Suma de los rangos de las diferencias positivas

T^- Suma de los rangos de las diferencias negativas

Empates

En ocasiones, los dos resultados de un par son iguales; entonces $x_i - y_i = d_i = 0$. En esta situación hay que excluir este tipo de par del análisis, disminuyendo por consiguiente el valor de "n".

El valor de la muestra n será el número total de pares objeto de estudio excepto los empates con $d = 0$.

Otro tipo de empate habitual es cuando dos o más diferencias son de la misma magnitud. En estas circunstancias se les asigna el mismo rango.

Investigación Comercial

El valor del rango se calcula de la forma siguiente:

Supongamos que tenemos tres pares cuyas diferencias son +1, -1 y -1; a cada par le asignaremos el rango 2; esto se debe a que promediamos los rangos que corresponden a cada diferencia. El correspondiente cálculo es:

$$(1 + 2 + 3) : 3 = 2$$

Al par siguiente le correspondería el rango 4, y así sucesivamente.

11.9.3.1 APLICACIÓN EN MUESTRAS PEQUEÑAS

En esta prueba, una muestra se considera pequeña cuando $n \leq 15$. En este caso, se aplica la tabla de rangos asignados de Wilcoxon que nos proporciona la probabilidad asociada a los valores T^+ . Si la probabilidad es menor o igual que el nivel de significación α , se rechaza H_0 .

11.9.3.2 MUESTRAS GRANDES

Una muestra se considera grande cuando n es mayor de 15. En este caso, se calcula la razón crítica Z , es decir, se compara con una normal (0,1) de media y desviación típica

$$m_T = \frac{n(n+1)}{4} \quad y \quad s_T = \sqrt{\frac{n(n+1)(2n+1)}{24}}$$

El valor de Z será:

$$Z = \frac{T^+ - m_T}{s_T}$$

Si existen rangos con empates, se corrige la varianza, utilizando la fórmula:

$$s_T^2 = \frac{n(n+1)(2n+1)}{24} - \frac{1}{2} \sum_{j=1}^g t_j(t_j-1)(t_j+1)$$

Donde g es el número de agrupamientos de diferentes rangos empatados y t_j número de rangos empatados agrupados en j .

Potencia eficacia.

Para muestras pequeñas es cercana al 95%.

Resumen de la prueba de Wilcoxon

- Para cada par se determina la diferencia y su signo $d = x_i - y_i$
- Se ordenan los rangos por valores absolutos (sin tener en cuenta el signo)
- A las diferencias que tengan el mismo valor se les asigna el rango promedio
- A cada rango se le asigna el signo + o - de la diferencia correspondiente
- Se determina el valor de n que es el número de diferencias distintas de 0
- Se determina T^+ , que es la suma de los rangos de signo positivo
- Se determina la significación en función del tamaño de n
- Si n es igual o menor de 15, la tabla de Wilcoxon nos proporciona la probabilidad asociada a los valores de T^+ . Si la probabilidad es igual o menor que el nivel de significación α , se rechaza la H_0
- Si n es mayor de 15, se calcula Z utilizando la fórmula

$$Z = \frac{T^+ - m_T}{s_T}$$

- En el caso de rangos con empate, se utiliza la correspondiente corrección para la desviación típica.
- La probabilidad asociada se determina con la tabla de la normal
- Para pruebas bidireccionales, se multiplica por dos el valor de tabla. Si la probabilidad obtenida de esta manera es menor o igual que α , se rechaza H_0

CASO PRÁCTICO

A un grupo de 10 consumidores potenciales se les pide que valoren, en una escala de 0 a 10 (0 como valor mínimo y 10 como máximo), dos refrescos A y B, respecto a un determinado atributo.

Los resultados obtenidos son los de la tabla siguiente:

Consumidor	1	2	3	4	5	6	7	8	9	10
Refresco A	7	9	6	5	8	5	7	9	7	4
Refresco B	6	7	4	9	7	6	7	6	9	9
Diferencia	1	2	2	-4	1	-1	0	3	-2	-5

Análisis de la Investigación Cuantitativa

Solución:

H_0 . La suma de rangos es nula

Hay 9 diferencias no nulas a las que les corresponden los siguientes rangos

		Rango
1, 1, -1	$(1 + 2 + 3) : 3 = 2$	2, 2 -2
2, 2, -2	$(4 + 5 + 6) : 3 = 5$	5, 5, -5
3	7	7
-4	8	-8
-5	9	-9

Los estadísticos correspondientes son:

$$T^+ = 2 + 2 + 5 + 5 + 7 = 21 \quad \text{y} \quad T^- = 2 + 5 + 8 + 9 = 24$$

El total de diferencias no nulas es $n = 9$

Buscando en la tabla T de Wilcoxon para un $\alpha = 0.05$ y $n = 9$ en la prueba bilateral obtenemos el valor de 6

Conclusión:

Como el valor de tablas es inferior al calculado $T^+ = 21$, no se puede rechazar la hipótesis nula. Esto significa que no hay diferencias significativas en la evaluación del atributo estudiado entre los dos refrescos.

Análisis de la Investigación Cuantitativa

TABLA T DE WILCOXON

a Test Unilateral			
N	0'025	0'01	0'05
a Test Bilateral			
N	0'05	0'02	0'01
6	0	-	-
7	2	0	-
8	4	2	0
9	6	3	2
10	8	5	3
11	11	7	5
12	14	10	7
13	17	13	10
14	21	16	13
15	25	20	16
16	30	24	20
17	35	28	23
18	40	33	28
19	46	38	32
20	52	43	38
21	59	49	43
22	66	56	49
23	73	62	55
24	81	69	61
25	89	77	68

11.10 CASO DE DOS MUESTRAS INDEPENDIENTES.

11.10.1 INTRODUCCIÓN

En la Investigación Comercial, en muchas ocasiones, no se pueden utilizar muestras relacionadas, utilizándose muestras independientes. En este tipo de investigación las dos muestras son obtenidas por uno de los siguientes procedimientos:

De forma aleatoria de dos poblaciones diferentes

De una misma población se elige una muestra aleatoria dentro de la cual se obtienen submuestras.

En ambos casos, no es necesario que el tamaño de las muestras sea idéntico.

En los test paramétricos, la prueba usual en el caso de muestras independientes es el test “t” a las medias de los dos grupos.

En los test no paramétricos, las pruebas más usuales son:

- Prueba exacta de Fisher para tablas de 2 x 2
- Prueba de Ji cuadrado para dos muestras independientes.
- Prueba de la mediana
- Prueba de Wilcoxon, Mann, Whitney
- Prueba de rangos ordenados (poderosa)
- Prueba de Kolmogorov, Smirnov para dos muestras
- Prueba de las permutaciones para dos muestras independientes
- Prueba de Siegel Tukey para diferencias en la escala
- Prueba de Moses para diferencias en la escala

Todas las pruebas no paramétricas para dos muestras independientes evalúan la hipótesis de que las dos muestras provienen de la misma población, las pruebas son más o menos sensibles a diferentes tipos de diferencias entre las muestras.

11.11 CASO DE K MUESTRAS RELACIONADAS

11.11.1 INTRODUCCIÓN

En determinados estudios de mercado interesa estudiar más de dos muestras simultáneamente, por ejemplo en experimentación. En estas circunstancias, es preciso

disponer de pruebas estadísticas que nos indiquen la posible diferencia global entre las “k” muestras. En las pruebas paramétricas se recurre al análisis de la varianza (test F).

Los test no paramétricos más utilizados son:

- Prueba Q de Cochran
- Análisis de varianza bifactorial, por rangos, de Friedman
- Prueba de Page para alternativas ordenadas

Estas pruebas son adecuadas cuando las mediciones de la variable están en escala ordinal.

11.12 CASO DE “K” MUESTRAS INDEPENDIENTES

11.12.1 INTRODUCCIÓN

Estas pruebas se utilizan cuando el investigador necesita decidir si varias muestras independientes pueden considerarse provenientes de la misma población.

La hipótesis nula a contrastar es que las k muestras independientes se han extraído de la misma población o de k poblaciones idénticas.

La prueba paramétrica habitual es el análisis de la varianza (test F).

Los test no paramétricos más usuales son:

- Test de Ji cuadrado para muestras independientes
- Prueba de la mediana (extensión)
- Análisis de varianza unifactorial por rangos de Kruskal, Wallis
- Prueba de Jonckheere para niveles ordenados de la variable.

12. BIBLIOGRAFÍA RECOMENDADA.

1. ANÁLISIS ESTADÍSTICO MULTIVARIABLE. Teoría y ejercicios. R. Sierra Bravo. Editorial Paraninfo, 1994
2. ANÁLISIS MULTIVARIANTE. (5ª edición). Hair, Anderson, Tatham, Black. Prentice may, 1999
3. APLICACIONES DE INVESTIGACIÓN COMERCIAL. Elena Abascal, Ildefonso Grande. ESIC editorial, 1994
4. ¿CÓMO HACER INVESTIGACIÓN DE MERCADOS? P. N. Hague, P.Jackson. Deusto, 1992
5. ¿CÓMO MEDIR LA SATISFACCIÓN DEL CLIENTE? Desarrollo y utilización de cuestionarios. Bob E. Hayes. Ediciones Gestión 2000, 1995
6. CUADERNOS DE ESTADÍSTICA: Análisis de varianza. Francisco J. Tejedor. Editorial La Muralla Hespérides, 1999
7. CUADERNOS DE ESTADÍSTICA: Análisis de correspondencias. Luis Joaristi Olariaga. Editorial La Muralla Hespérides, 1999
8. CUADERNOS DE ESTADÍSTICA: Análisis factorial. E. García Jiménez y otros. Editorial La Muralla Hespérides, 1999
9. CUADERNOS DE ESTADÍSTICA: El análisis multivariante en la investigación científica. Rosario Martínez Arias. Editorial La Muralla Hespérides, 1999
10. CUADERNOS DE ESTADÍSTICA: Regresión múltiple. Juan Etxeberría. Editorial La Muralla Hespérides, 1999
11. CUADERNOS METODOLÓGICOS: Cuestionarios. (26) Mª José Azofra. CIS Centro de Investigaciones Sociológicas, 1999
12. DIRECTORIO DE FUENTES DE INFORMACIÓN DE LA ECONOMÍA ESPAÑOLA. Paloma Portela. Directora. Crítica, 1996
13. DISEÑO DE INVESTIGACIONES: Cuaderno de Prácticas. Hilda Gambara. Mc Graw Hill, 1995
14. DISEÑO DE INVESTIGACIONES: Introducción a la lógica de la investigación en Psicología y Educación. Orfelio G. León, Ignacio Montero. Mc Graw Hill, 1993

15. DISEÑO Y TRATAMIENTO ESTADÍSTICO DE ENCUESTAS PARA ESTUDIOS DE MERCADO. Julián Santos Peñas y otros. Editorial Centro de Estudios Ramón Areces SA, 1999
16. DYANE Versión 2 Diseño y análisis de encuestas en investigación social y de mercados. Miguel Santesmases Mestre. Pirámide, 2001
17. EL ABC DE INTERNET y las 1000 direcciones más útiles. Esine, 2000
18. EL ARTE DE LA ENCUESTA. Principios básicos para no especialistas. Y. Harvatopoulos y otros. Deusto, 1992
19. EL DIFERENCIAL SEMÁNTICO. Técnicas de investigación social. Alfredo Bechini Tejados. Hispano Europea, 1986
20. EL MÉTODO DELPHI. Una técnica de previsión para la incertidumbre. Jon Landeta. Ariel Practicum, 1999
21. EL SONDEO UNA HERRAMIENTA DE MARKETING. J. Antoine. Deusto, 1992
22. ESTADÍSTICA APLICADA.(2ª edición) Félix Calvo. Ediciones Deusto, 1994
23. ESTADÍSTICA NO PARAMÉTRICA Aplicada a las ciencias de la conducta. Sydney Siegel N. John Castellan. Trillas,1995
24. FUNDAMENTOS Y TÉCNICAS DE INVESTIGACIÓN COMERCIAL. (5ª edición) Ildefonso Grande Esteban, Elena Abascal Fernández. Editorial ESIC, 2000
25. IDENTIFICACIÓN DE LOS MERCADOS APROPIADOS. David Parmelee. Granica, 1998
26. INTRODUCCIÓN AL ANÁLISIS ECONÓMETRICO CON DATOS DE PANEL. Manuel Arellano. Servicio de estudios Banco de España.
27. INVESTIGACIÓN COMERCIAL DINÁMICA. Luis Roig Sancho. Deusto, 1982
28. INVESTIGACIÓN COMERCIAL, 22 casos prácticos y un apéndice teórico. Mª Ángeles González Lobo. Esic, 2000
29. INVESTIGACIÓN DE MARKETING. Teodoro Luque. Ariel, 1997
30. INVESTIGACIÓN DE MERCADOS (3ª edición). David Aaker, Gerge S. Day. Mc Graw Hill, 1994

31. INVESTIGACIÓN DE MERCADOS (5ª edición). Kinneer Taylor. Mc Gaw Hill, 1998
32. INVESTIGACIÓN DE MERCADOS (6ª edición). William G. Zikmund. Prentice Hall, 1998
33. INVESTIGACIÓN DE MERCADOS Ronald M. Weiers. Prentice may, 1986
34. INVESTIGACIÓN DE MERCADOS Y ESTRATEGIA DE MARKETING. Laurentino Bello y otros. Editorial Cívitas,1993
35. INVESTIGACIÓN DE MERCADOS. Cómo se realiza, cómo se utiliza. Ramón Ribas Muntan. Editorial Index, 1993
36. INVESTIGACIÓN DE MERCADOS. Guía Maestra para el profesional. Jeffrey Pope. Editorial Norma, 1981
37. INVESTIGACIÓN DE MERCADOS. Jeffrey Pope. Parramón, 1994
38. INVESTIGACIÓN DE MERCADOS. Salvador Miquel, y otros. Mc Graw Hill, 1996
39. INVESTIGACIÓN DE MERCADOS.Un enfoque práctico. Narres K. Malhotra. Prentice Hill, 1997 (2ª edición)
40. INVESTIGACIÓN DE MERCADOS: Obtención de información. Ángel Fernández Nogales. Editorial Cívitas, 1997
41. INVESTIGACIÓN EN MARKETING. Enrique Díez de Castro, Javier Landa Bercebal. Editorial Cívitas, 1994
42. INVESTIGACIÓN INTEGRAL DE MERCADOS. Un enfoque operativo. José Nicolás Jany. Mc. Graw Hill, 1994
43. INVESTIGACIÓN INTEGRAL DE MERCADOS. Un enfoque para el siglo XXI. José Nicolás Jany. Mc Graw Hill, 2000
44. INVESTIGACIÓN Y ANÁLISIS DE MERCADO. Lehmann. CECSA, 1993.
45. LA INTEGRACIÓN DE LOS MÉTODOS CUANTITATIVO Y CUALITATIVO EN LA INVESTIGACIÓN SOCIAL. Significado y medida. Eduardo Bericat. Ariel Sociología, 1998
46. LA INVESTIGACIÓN CIENTÍFICA DE LOS MEDIOS DE COMUNICACIÓN. Una introducción a sus métodos. Roger D. Wimmer, Joseph R. Dominick. Bosch, 1996

47. LA INVESTIGACIÓN COMERCIAL COMO SOPORTE DEL MARKETING. Ramón Pedret. Deusto, 2000
48. LA INVESTIGACIÓN EN MARKETING. (Dos tomos) Varios Autores. Aedemo, 2000
49. LA INVESTIGACIÓN EN RELACIONES PÚBLICAS. John V. Pavlik. Gestión 2000, 1999
50. LA PRÁCTICA DE LA INVESTIGACIÓN COMERCIAL. Francisco Serrano Gómez. Editorial ESIC, 1990
51. LOS ESTUDIOS DE MERCADO. José M^a Ferré Trenzano. Jordi Ferré Nadal. Díaz de Santos, 1997.
52. MANUAL DE INVESTIGACIÓN COMERCIAL. Enrique Ortega Martínez. Pirámide, 1990
53. MANUAL DE PSICOLOGÍA EXPERIMENTAL. Metodología de investigación. Juan Pascual y otros. Ariel, 1996
54. MANUAL PARA ENCUESTADORES. V. G. Manzano y otros. Ariel, 1996
55. MEDICIÓN, INVESTIGACIÓN E INFORMACIÓN DE LA PUBLICIDAD. Raúl Eguizabal, Antonio Caro. Comunicación, 2000
56. METODOLOGÍA DE LA INVESTIGACIÓN PARA ADMINISTRACIÓN Y ECONOMÍA. César Augusto Bernal. Prentice Hall, 2000
57. METODOLOGÍA DE LA OBSERVACIÓN EN LAS CIENCIAS HUMANAS. M^a Teresa Anguera. Cátedra, 1992 (5^a edición)
58. METODOLOGÍA PARA LA INVESTIGACIÓN EN MARKETING Y DIRECCIÓN DE EMPRESAS. Francisco José Sarabia Sánchez (Coordinador), Pirámide, 1999
59. MÉTODOS MULTIVARIANTES PARA LA INVESTIGACIÓN COMERCIAL. Elena Abascal, Ildefonso Grande. Ariel Economía, 1989
60. MODELOS CAUSALES. Técnicas de investigación social. B. Visauta Vinacua. Hispano Europea, 1986
61. PLANEACIÓN PROSPECTIVA. Una estrategia para el diseño del futuro. Tomás Miklos, M^a Elena Tello. Noriega Limusa, 1991
62. PREPARACIÓN, TABULACIÓN Y ANÁLISIS DE ENCUESTAS PARA DIRECTIVOS. Joseph Múria Albiol y otros. ESIC, 1998

63. ¿QUÉ ES LA INVESTIGACIÓN DE MERCADOS? Jack Hamilton. AEDEMO ESOMAR, 1989
64. TÉCNICAS DE ANÁLISIS DE DATOS EN INVESTIGACIÓN DE MERCADOS. Teodoro Luque Martínez (Coordinador). Pirámide, 2000
65. TÉCNICAS DE INVESTIGACIÓN APLICADAS A LAS CIENCIAS SOCIALES. Jorge Padua El Colegio de México. Fondo de Cultura, 1992
66. TÉCNICAS DE INVESTIGACIÓN SOCIAL. Teoría y ejercicios. R. Sierra Bravo Paraninfo, 1991
67. TÉCNICAS DE LA INVESTIGACIÓN SOCIAL. Fernando Giobellina Brumana. Nueva Escuela Publicaciones, 1995
68. TÉCNICAS ESTADÍSTICAS CON SPSS. César Pérez. Prentice Hall, 2001
69. TEMAS DE INVESTIGACIÓN DE MEDIOS PUBLICITARIOS. J. Enrique Bigné. Editorial ESIC, 2000

Zaragoza a 20 de Agosto de 2004